



Universiteit  
Leiden  
The Netherlands

## **Taxonomic classification and abundance estimation using 16S and WGS-A comparison using controlled reference samples**

Khachatryan, L.; Leeuw, R.H. de; Kraakman, M.E.M.; Pappas, N.; Raa, M. te; Mei, H.L.; ... ; Laros, J.F.J.

### **Citation**

Khachatryan, L., Leeuw, R. H. de, Kraakman, M. E. M., Pappas, N., Raa, M. te, Mei, H. L., ... Laros, J. F. J. (2020). Taxonomic classification and abundance estimation using 16S and WGS-A comparison using controlled reference samples. *Forensic Science International: Genetics*, 46. doi:10.1016/j.fsigen.2020.102257

Version: Publisher's Version

License: [Creative Commons CC BY-NC-ND 4.0 license](https://creativecommons.org/licenses/by-nc-nd/4.0/)

Downloaded from: <https://hdl.handle.net/1887/3181651>

**Note:** To cite this publication please use the final published version (if applicable).



## Short communication

## Taxonomic classification and abundance estimation using 16S and WGS—A comparison using controlled reference samples

Lusine Khachatryan<sup>a,\*</sup>, Rick H. de Leeuw<sup>a</sup>, Margriet E.M. Kraakman<sup>b</sup>, Nikos Pappas<sup>c</sup>, Marije te Raa<sup>a</sup>, Hailiang Mei<sup>c</sup>, Peter de Knijff<sup>a</sup>, Jeroen F.J. Laros<sup>a,d</sup>

<sup>a</sup> Department of Human Genetics, Leiden University Medical Center, Leiden, the Netherlands

<sup>b</sup> Department of Medical Microbiology, Leiden University Medical Center, Leiden, the Netherlands

<sup>c</sup> Sequencing Analysis Support Core, Leiden University Medical Center, Leiden, the Netherlands

<sup>d</sup> Department of Clinical Genetics, Leiden University Medical Center, Leiden, the Netherlands

## ARTICLE INFO

## Keywords:

Metagenomics  
Skin microbiome  
Taxonomic profiling  
16S sequencing  
WGS sequencing

## ABSTRACT

The assessment of microbiome biodiversity is the most common application of metagenomics. While 16S sequencing remains standard procedure for taxonomic profiling of metagenomic data, a growing number of studies have clearly demonstrated biases associated with this method. By using Whole Genome Shotgun sequencing (WGS) metagenomics, most of the known restrictions associated with 16S data are alleviated. However, due to the computationally intensive data analyses and higher sequencing costs, WGS based metagenomics remains a less popular option. Selecting the experiment type that provides a comprehensive, yet manageable amount of information is a challenge encountered in many metagenomics studies. In this work, we created a series of artificial bacterial mixes, each with a different distribution of skin-associated microbial species. These mixes were used to estimate the resolution of two different metagenomic experiments - 16S and WGS - and to evaluate several different bioinformatics approaches for taxonomic read classification. In all test cases, WGS approaches provide much more accurate results, in terms of taxa prediction and abundance estimation, in comparison to those of 16S. Furthermore, we demonstrate that a 16S dataset, analysed using different state of the art techniques and reference databases, can produce widely different results. In light of the fact that most forensic metagenomic analysis are still performed using 16S data, our results are especially important.

## 1. Introduction

In recent years, metagenomics - the genomic analysis of microorganisms by direct extraction of DNA from an environmental sample - has become one of the most rapidly developing branches of microbiology [1–3]. The interest in metagenomics has grown drastically due to the expanding number of studies showing that the vast majority of microorganisms cannot be grown under laboratory conditions [4–7]. The possibility of culture-free investigation of microbial biodiversity directly from an environmental habitat led to many studies benefiting a wide range of fields such as human health [8–12], ecology [13,14], agriculture [15–17], forensics [18,19], food and drugs production [20–22]. Taxonomic profiling of metagenomic data is a key step during the data analysis, allowing researchers to understand the structure of a microbiome and to estimate relative abundances of the organisms living in it. The main goal of this study is to compare different data types and methods for taxonomic profiling of metagenomic data sets with known

abundance distributions of inhabitants.

The most common technique to investigate microbiome composition is amplicon-based sequencing of the 16S rRNA gene [23,24]. This relatively short (~1500 bp) gene is universal among bacteria and archaea [25,26]. There are in total nine hypervariable regions in the 16S rRNA gene that provide phylogenetic signatures on different taxonomic levels. Hypervariable regions are surrounded by highly conserved sequences, which are used for primer design. The analysis of 16S metagenomic datasets is usually performed in combination with one of several curated databases that contain annotated sequences of the 16S rRNA gene or its parts [27]. The most commonly used 16S-specific databases are RDB [28,29], GreenGenes [30] and SILVA [31]. Analysis of 16S data is now routine for metagenomic-associated projects, though many studies demonstrated a number of biases associated with this type of data that make the validity of this approach questionable. Several reports stressed uneven coverage of microorganisms' diversity spectrum by common PCR primers for the 16S rRNA gene amplification [32–37].

\* Corresponding author.

E-mail address: [l.khachatryan@lumc.nl](mailto:l.khachatryan@lumc.nl) (L. Khachatryan).

<https://doi.org/10.1016/j.fsigen.2020.102257>

Received 9 June 2019; Received in revised form 30 December 2019; Accepted 27 January 2020

Available online 05 February 2020

1872-4973/ © 2020 The Authors. Published by Elsevier B.V. This is an open access article under the CC BY-NC-ND license (<http://creativecommons.org/licenses/by-nc-nd/4.0/>).

Second, the 16S rRNA gene does not have a correct phylogenetic relationship within particular taxa [38,39]. The fact that bacteria and archaea might carry different copy numbers of the 16S rRNA gene in their genomes seriously influences a reliable abundance estimation after analysis of 16S data [40]. Additionally, the choice of a specific hypervariable region and the reference database for the subsequent analysis requires *a priori* knowledge about the investigated metagenome. Lastly, 16S data cannot be used to investigate the metagenome functional profile, nor does it provide any information about eukaryotic or viral members of the microbial community. The applicability of 16S data was shown for a set of forensic studies. For example, 16S data was successfully used for body fluid recognition [41] or matching between individuals' skin datasets and touched objects [42,43]. The success of such analyses, however, does not imply that a 16S-based analysis of all metagenomic data is reliable (or possible).

Apart from 16S, there are other methods that use rRNA genes to investigate microbial diversity. Among them are 23S, 5S, 12S and various combinations [44–46]. Other methods like the IS-pro approach use 16S-23S ribosomal interspace fragment lengths to analyse microbial communities [47]. Although these methods are very suitable for some specific tasks, they are not as widely applied as 16S. Several recent studies are based on targeting other genes in addition to 16S in order to determine the cell type of the forensic traces [48] or to perform skin sample identification using only microbial targeting genes [49,50]. These studies also suggest that traditional 16S data is not always sufficient for a meaningful metagenomic analysis of forensic traces.

In recent years, the number of metagenomic studies based on the whole genome shotgun (WGS) sequencing data type has grown [51–55]. Among the main reasons for this are advantages in sequencing techniques allowing for the generation of sufficient number of high-quality reads for the WGS datasets, and bioinformatics algorithms to perform subsequent analysis of the big data. Though using WGS data avoids the biases introduced by 16S, it requires more computationally intense analysis, as well as higher sequencing costs.

While many studies in the field of forensics are based on the analysis of 16S data [56], “the capacity of WGS data of microbiomes to aid in forensic investigations by connecting objects and environments to individuals has been poorly investigated” [57]. Presently, WGS experiments are reserved for those studies for which analysis beyond the taxonomical assignment is required: investigating the microbiomes' functional profile, correlation between metagenome and host genome, search for the possible virulent genes, etc. The vast majority of taxonomical annotations is still performed by using only 16S data, despite all known disadvantages of the method [51]. One of the reasons for that is the lack of a well-performed benchmark study, comparing 16S and WGS data types. The vast majority of existing metagenomics benchmarks are created in order to evaluate the accuracy of various metagenomic profiles and comprise either only 16S [58] or only WGS data [59–66]. Existing benchmarks that can be used to compare 16S and WGS data types are *in-silico* created and based on a random set of bacterial species, lacking the information about whether or not the selected set of organisms might live together in the same environment [67]. One of the main goals of this study is the creation of a set of benchmarks allowing to compare the 16S and WGS data types using a set of *in-vitro* DNA mixes of bacteria species inhabiting skin.

Over the last decade, the number of different techniques for metagenomics data analysis has grown remarkably. The tools used for performing the taxonomical annotation, can be split into several groups based on the following criteria: strategy for reads assignment (alignment or matching based on the *k*-mers or sequences signatures); the database against which the search is performed; the proportion of reads participating in the profiling (all reads, only one read per read group, only reads with particular features).

To investigate which type of metagenomic data is preferable for accurate taxonomic annotation, as well as to test which method of reads assignment yields more precise output, we created a series of bacterial

mixes with known content. Each metagenomic mix incorporated 14–15 bacterial species belonging to 7 distinct bacterial genera. Each mix had a distinct distribution of the species abundances. For the analysis we selected two popular tools: Centrifuge [68] and MG-RAST [69]. These allow analysis of amplicon and WGS sequencing data and both perform the metagenome profiling by a comparison of sequencing data to a reference database. However, the strategies for metagenome profiling they exploit are different.

We did not include other popular tools for metagenomic analysis in this study as they either have a similar analysis strategy as the tools described above or are designed only for WGS or amplicon data analysis. In many studies, QIIME [70], objectively the most popular tool for amplicon data analysis, was shown to perform with the same accuracy as the MG-RAST pipeline for 16S rRNA sequencing data [71].

## 2. Materials and methods

### 2.1. DNA extraction and concentration measurement

Laboratory pure cultures of 15 bacterial species that frequently inhabit human skin (Table 2) were grown with gentle shaking overnight at 37°C. Genomic DNA was isolated with the Easy-DNA™ gDNA Purification Kit (Invitrogen™ Thermo Fisher Scientific) using the standard protocol with ethanol precipitation [72]. RNA contamination was removed using RNase A (Roche) and the DNA was stored at 4°C. DNA concentrations were measured with the Qubit 3.1 Fluorometer (Invitrogen™).

### 2.2. Metagenomic mixes creation

Four bacterial mixes with known genome abundances were created for this research. In order to achieve the desired species abundances, the estimated genome size and the measured DNA concentration for each bacteria were used. One mix was created to have a uniform- and other three mixes an exponential ( $\lambda = 1/6$ ,  $\lambda = 1/2$  and  $\lambda = 5/6$ ) distribution of species abundances. From here on, these mixes are referred to as EQ, EXP16, EXP12 and EXP56 respectively. Due to technical reasons, *Corynebacterium jeikeium* was included only in EQ. The remaining 14 species were used in all mixes.

### 2.3. WGS sequencing library creation

DNA shearing was performed using the Covaris S2 sonicator (Covaris®) with the following settings: duty factor = 10 %, intensity = 2.5, cycles/burst = 200, temperature = 6°C, total time, sec = 45. Size selection was performed on the sheared products with Ampure XP beads (Agencourt) to maintain insert size around 450 base pairs.

Illumina sequencing libraries were prepared by ligating custom Illumina Truseq adapters with dual barcoding (10 base pairs) using the KAPA Hyper Prep Library Preparation kit (KAPA Biosystems, Inc.). To increase library yield, additional library amplification was performed with KAPA HIFI HotStart ReadyMix using the PCR protocol described in Table 1. To enable balanced pooling, sequencing libraries were quantified in duplicate by real time PCR using the KAPA SYBR® FAST qPCR kit. Quantification reactions were performed on a LightCycler® 480

**Table 1**  
PCR protocol for the WGS library preparation.

| Step                 | Temperature, °C | Duration, min | Cycles              |
|----------------------|-----------------|---------------|---------------------|
| Initial denaturation | 95              | 3             | 1, hold             |
| Denaturation         | 98              | 0.25          | Ranged from 3 to 8  |
| Annealing            | 59              | 0.5           | depending on sample |
| Extension            | 72              | 1.5           |                     |
| Final extension      | 72              | 5             | 1, hold             |

**Table 2**  
Bacterial species used for metagenomics mixes.

| Bacteria                                       | Number of contigs | Accession number | Reference length, Mb | Total assembly length, Mb |
|------------------------------------------------|-------------------|------------------|----------------------|---------------------------|
| <i>Acinetobacter johnsonii</i> ATCC 17969      | 206               | NZ_CP010350.1    | 3.51                 | 3.88                      |
| <i>Acinetobacter lwoffii</i> ATCC 15309        | 180               | NA               | NA                   | 3.44                      |
| <i>Corynebacterium jeikeium</i> ATCC 43734     | 234               | NC_007164.1      | 2.46                 | 2.6                       |
| <i>Corynebacterium urealyticum</i> ATCC 43042  | 99                | NC_010545.1      | 2.37                 | 2.35                      |
| <i>Moraxella osloensis</i> NCTC 10145          | 89                | CP014234.1       | 2.43                 | 2.58                      |
| <i>Propionibacterium acnes</i> ATCC 6919       | 26                | NC_017550.1      | 2.49                 | 2.55                      |
| <i>Pseudomonas aeruginosa</i> ATCC 10145       | 99                | NC_002516.2      | 6.26                 | 6.35                      |
| <i>Staphylococcus aureus</i> ATCC 29213        | 45                | NZ_CP009361.1    | 2.78                 | 2.72                      |
| <i>Staphylococcus capitis</i> ATCC 27840       | 52                | NZ_CP007601.1    | 2.44                 | 2.6                       |
| <i>Staphylococcus epidermidis</i> ATCC 12228   | 142               | NC_004461.1      | 2.5                  | 3.3                       |
| <i>Staphylococcus haemolyticus</i> ATCC 29970  | 770               | NC_007168.1      | 2.69                 | 2.86                      |
| <i>Staphylococcus saprophyticus</i> ATCC 15305 | 351               | NC_007350.1      | 2.15                 | 1.89                      |
| <i>Streptococcus pyogenes</i> ATCC 19615       | 65                | NZ_CP008926.1    | 1.84                 | 1.82                      |
| <i>Staphylococcus xylosus</i> ATCC 29971       | 97                | NZ_CP008724.1    | 2.52                 | 2.74                      |
| <i>Streptococcus mitis</i> LMG 14552           | 49                | NC_013853.1      | 2.76                 | 2.83                      |

(Roche) using a dilution series of PhiX control library (Illumina) as standard [72]. After pooling the libraries, the final pool was quantified again using the same method to enable optimal loading of the flow cell.

#### 2.4. 16S sequencing library creation

Previously published [73] Primers and PCR-protocol for the amplification of V3-V4 region of the 16S rRNA were used. Illumina sequencing libraries were prepared by ligating custom Illumina Truseq adapters with dual barcoding (10 base pairs) using the KAPA Hyper Prep Library Preparation kit (KAPA Biosystems, Inc.).

#### 2.5. DNA sequencing

Sequencing of WGS and 16S libraries was performed on the MiSeq®sequencer (Illumina) using v3 sequencing reagents according to the manufacturer's protocol with approximately 5 % of PhiX control. This yielded one paired-end dataset with a read length of 299bp per sample.

#### 2.6. Bacterial genomes assembly

Sequencing reads for each bacterium were preprocessed using the Flexiprep quality control pipeline [74].

Post-QC reads were assembled by SPAdes Genome Assembler [75] with default settings.

#### 2.7. Regression analysis

*k*-mer counting was performed using command *count* of the kPAL toolkit [76] with *k* set to 11. In case of the absence of the alternative DNA stand, *k*-mer profiles were balanced with *balance* command of the kPAL toolkit. Linear regression was done using the *sklearn.linear\_model* class of the scikit-learn package for Python [77] with the *fit\_intercept* parameter set to "False". The model training and prediction was performed using 5-fold Cross Validation.

#### 2.8. Analysis using Centrifuge

Centrifuge is a popular tool that enables a fast classification of reads in a metagenomic dataset using comparison of *k*-mers derived from each read to an indexed database. Centrifuge performs classification for all reads in a metagenomic dataset independently using the following algorithm. A fast and effective comparison is achieved using the genome indexing technique, which is based on the Burrows-Wheeler transform [78] and the Ferragina-Manzini index [79]. To perform taxonomy assignment, Centrifuge requires an indexed database which

is based on the reference database and its associated phylogenetic tree. A number of popular and regularly updated premade indexed databases are available on the Centrifuge website [80]. It is also possible to create a custom Centrifuge indexed database.

Metagenomic mixes datasets were subjected to a QC-check using FastQC (version 0.11.7, [81]). Leftover adapter removal and quality trimming of the reads was performed with cutadapt [82] (version 1.16, using options *-trim-n*, *-minimum-length* = 50 and *-quality-cutoff* = 20). The number of reads before and after each aforementioned step can be found in Supplementary Table S1. High quality pairs of overlapping reads were merged with FLASH [83] (version 1.2.11, using option *-max-overlap* = 300). For the subsequent taxonomic classification with Centrifuge, both merged reads and pairs of non-merged reads were used.

Post-QC reads were analysed with Centrifuge (version 1-0-2-beta, default settings). Three different reference databases were used for the analysis: RefSeq database of complete genomes of bacteria and archaea [84] (downloaded as premade in April 2018 Centrifuge index); GreenGenes 16S sequences database (downloaded in June 2018) and SILVA 16S sequences database (version SSURf Nr99, downloaded June 2018). In order to make the content of reference databases comparable, sequences marked as eukaryotic were removed from SILVA database. Results obtained by Centrifuge were analysed using the Pavian interactive browser [85] application.

#### 2.9. Analysis using MG-RAST

MG-RAST is a web-based tool that allows the user to upload sequences and their metadata and download the analysis results. The MG-RAST pipeline creates a metagenomic profile by extracting rRNA and protein coding sequences. Gene calling is performed by the FragGeneScan [86] algorithm, predicted protein sequences are clustered using UCLUST [87]. Potential rRNA genes are identified using BLAT [88] against a reduced version of the SILVA database and clustered with UCLUST. From each obtained cluster one representative sequence (the longest one) is chosen for the comparison with a reference database (M5nr58 [89] for proteins and combination of SILVA59, GreenGenes42 and RDP41 for rRNA analysis) using BLAT. All sequences from a particular cluster are assigned to the same taxonomic group as the clusters' representative. Thus, only rRNA genes and functional genes are used for the analysis of the metagenome, and the reads assignments are not independent. This strategy allows MG-RAST to perform taxonomic and functional profiling of metagenomic data. Finally, MG-RAST supports different metagenomics datatypes: genomic (including WGS and 16S) and transcriptomic. It also considers the metagenome origin, sequencing platform and many other features to tune the pipeline for a specific task.

Raw reads of bacterial mixes datasets were submitted to the MG-

RAST Metagenomics analysis server under project number 85582. Paired reads merging and quality control was performed as part of the standard MG-RAST pipeline.

### 2.10. Taxa abundance estimation and results evaluation

Since the 16S amplification product has the same length among all bacterial taxa, no correction for genome length is needed when estimating relative abundances of the taxa. For the WGS datasets however, normalization of read counts is required because of the differences in genome lengths. In order to perform correct taxa abundance estimations for taxonomic ranks higher than species, it is important to know how many reads assigned to that taxon belong to each species within the taxon. Both tools, Centrifuge and MG-RAST, assign reads to a node in the phylogenetic tree. Thus, reads assigned to a particular genus, for example, might belong to each of the species included to that genus as well as to the genus itself. The main assumption of our approach for the estimation of taxa abundances is the following: All reads, assigned to the node higher than species level (regardless of whether or not they have species annotation), will be distributed among the species belonging to that node the same way as the reads with known species annotation. If the estimated abundances for species were known (in case of taxonomic annotation with Centrifuge), the procedure is trivial. When performing the analysis with MG-RAST the reads are classified only up to the genus level. In that case an equal distribution of reads among the species belonging to the particular genus was assumed.

### 2.11. Statistical and correlation analysis

Correlation analysis was performed using the Pearson correlation coefficient, pair wise comparisons were performed using the two-sided Mann–Whitney *U* test [90] and False Discovery Rate (FDR, a statistical approach used in multiple hypothesis testing to correct for multiple comparisons) control was performed using the Benjamini–Hochberg procedure [91]. We used the ratio of properly predicted taxa to all taxa predicted at that rank as a measure for the precision. Sensitivity was calculated as the ratio of properly predicted taxa to all taxa that were

supposed to be present in the sample at that rank. F-scores (a measure of accuracy that considers both precision and sensitivity) were calculated as described in [92].

## 3. Results and discussions

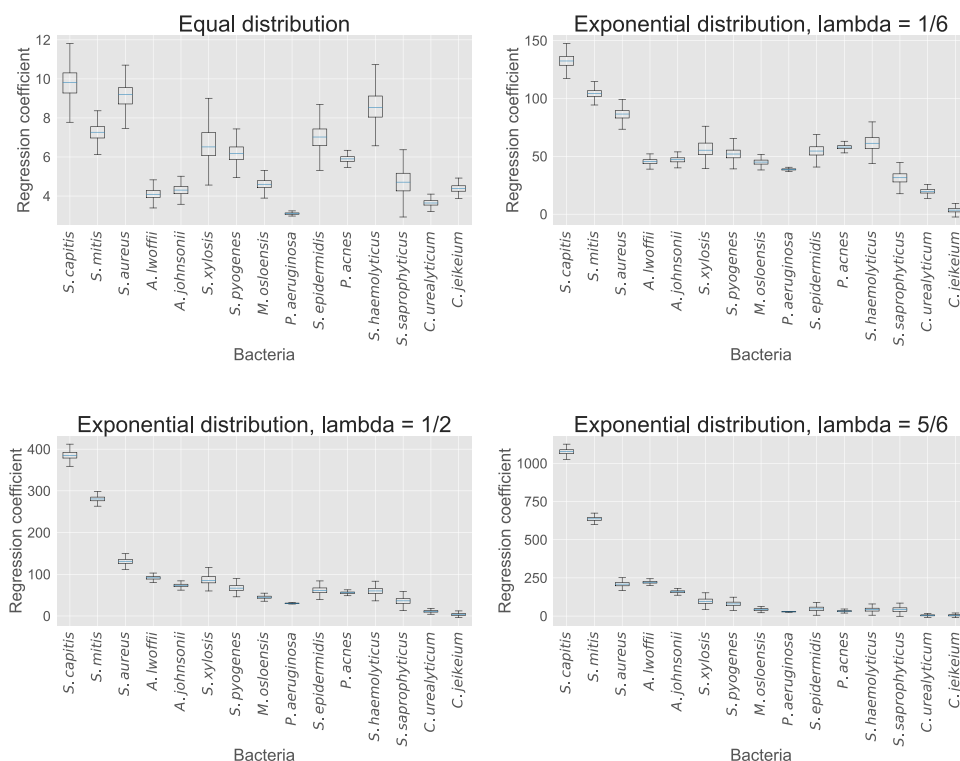
### 3.1. Individual bacterial genomes assembly

We sequenced and assembled the genomes of all 15 selected skin-associated bacteria individually. The total length of the assembly for each species was comparable to the length of the species references (see Table 2 and Section S1 of the Supplementary materials).

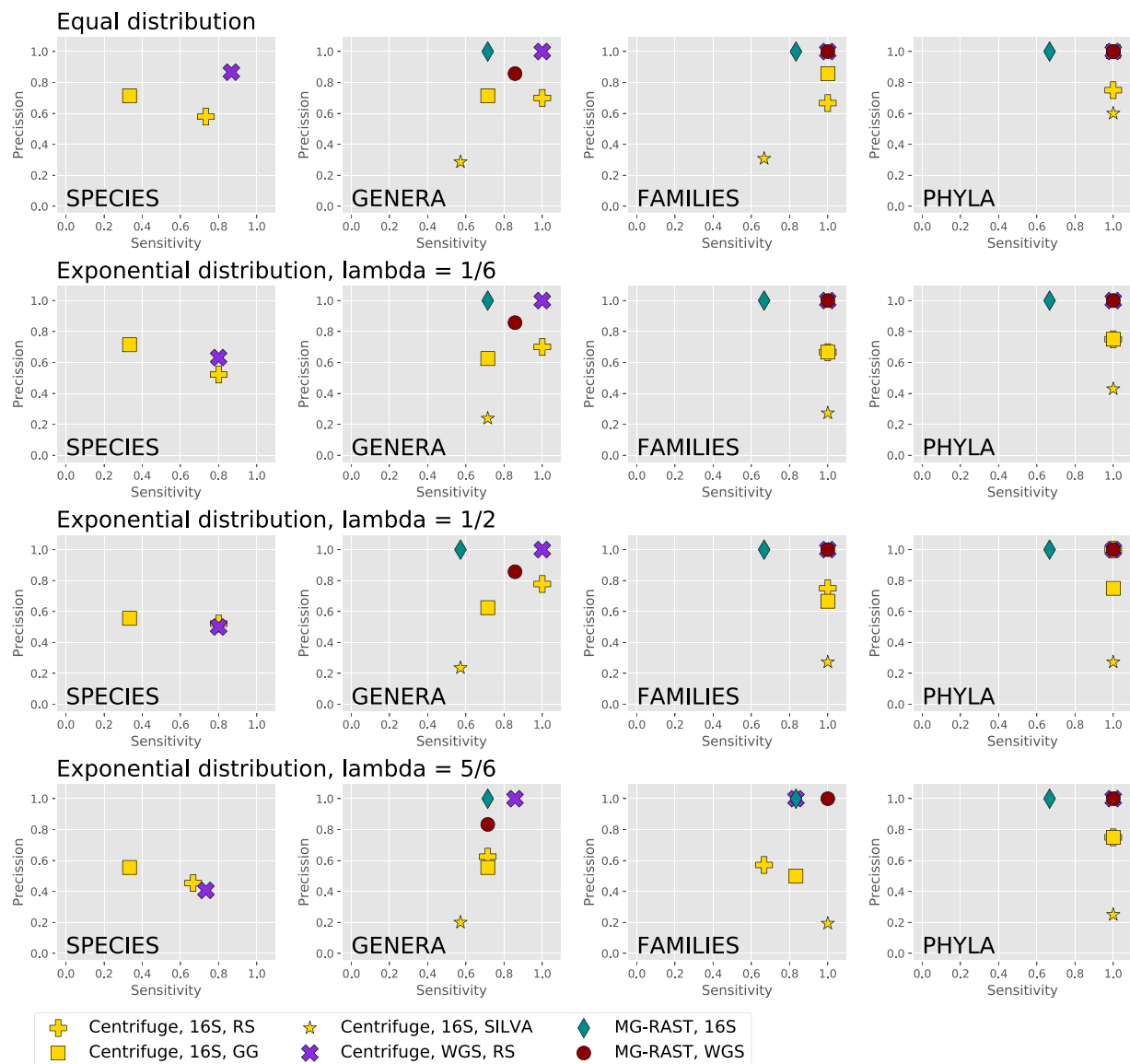
For one species (*Acinetobacter lwoffii*) there was no reference sequence available. Obtained assembly lengths as well as the DNA concentration measured for each bacterium were used to create four metagenomic mixes: one with equal and three with exponential ( $\lambda = 1/6$ ,  $\lambda = 1/2$  and  $\lambda = 5/6$ ) distribution of bacterial species abundances. Taxa abundances were ordered from high to low as shown in Fig. 1.

### 3.2. Estimation of reference abundances

In order to estimate an abundance of an organism in terms of genome copies, the length of the genome and the lengths and (relative) copy numbers of any plasmids needs to be known. In the absence of a strain-specific reference sequence, a *de novo* assembly of a single organism can be used to obtain these data [93]. In most common approaches [94], the coverage (and thereby the copy number) of contigs (see Supplementary Fig. S1) is not considered when estimating an assembly length, which leads to an inaccurate estimation of the organisms' genome length and thus influence the accuracy when creating bacterial mixes (see Supplementary Fig. S2 for a step-by-step explanation). Other factors, such as inaccuracy in DNA concentration measurement or mixing, can also lead to different abundances in the final bacterial mixes from those intended. Since the content of all our metagenomic mixes is known and individual assemblies of all bacterial species were available, the intended distribution of bacterial abundances in the metagenomic mixes could be verified using the following



**Fig. 1.** Regression analysis performed for metagenomic mixes to estimate relative abundances. Results for each mix are shown in a separate plot. Each boxplot represents the distribution of regression coefficients (vertical axes) obtained for a particular organism (horizontal axes), thus representing its relative abundance within that particular mix.



**Fig. 2.** Precision vs. sensitivity of different profiling approaches. Results for each mix and taxonomic rank are shown separately, with the sensitivity on the horizontal axes and the precision on the vertical axes. Each shape represents a combination of method, data type and reference database. RS – RefSeq database, GG – GreenGenes database, SILVA – SILVA database.

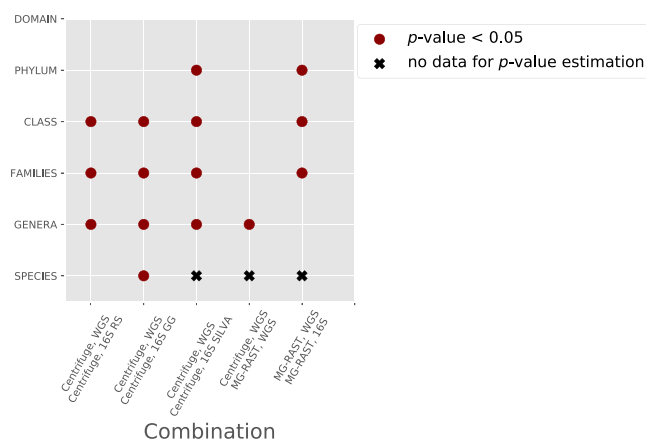
approach. We used  $k$ -mer counts as a proxy for the number of genomes present in a pure (unmixed) sample. Using these counts, we are able to infer the relative contributions to a mixture. We use randomly chosen  $k$ -mers from the pure samples as profiles for the organisms, the same  $k$ -mers are used to make a profile of the mix and by linear regression, we estimate the contribution of each profile and thereby the contribution of each organism to the mix. For a more detailed description and a motivational example, see Section S1 and Fig. S2 of the Supplementary materials. We calculated the 11-mer profiles for each bacteria using the contigs obtained after individual genome sequencing and assembly. Since profiles were calculated using contigs, we compensated for the absence of the reverse-complement DNA strand. We also calculated the 11-mer profiles of the WGS datasets of each of the metagenomic mixes, in these cases strand balancing was not applied. The 11-mer profiles were used to build a linear regression model in which the individual bacterial  $k$ -mer counts were treated as independent variables and the  $k$ -mer counts of the metagenomic mix served as dependent variable.

Since  $k$ -mer counts within one profile might be correlated, which violates the condition for using the regression analysis, we did not analyse the complete profile of 4,194,304 possible 11-mers. Instead we

performed 1000 iterations, in each iteration choosing 10,000 random  $k$ -mers and performing the regression analysis on that subset of  $k$ -mers. Thus, for each organism we got 1000 estimations of its abundance in each mix. The result of this analysis is presented in Fig. 1. Each boxplot shows the distribution of the organisms' abundances obtained from the regression analysis. The median model fit of the cross-validated models (measured using the  $R^2$  coefficient of determination) for each mix was larger than 0.95, accuracy of the prediction (also measured using the  $R^2$  but on the data that did not participate in the model training) ranged from 0.8 to 0.92 depending on the mix.

The regression analysis confirmed the distribution of bacterial abundances we aimed for (uniform distribution turning into the exponential one), though for some species (e.g., *S. haemoliticus* and *P. aeruginosa*), slight positive or negative deviations from the anticipated values were found. This can be caused by a number of factors such as inaccuracy in the DNA concentration measurement or DNA mixing, presence of large amounts of non-chromosomal DNA (e.g., plasmids) in the pool of bacterial DNA or inaccuracy in bacterial genome size estimation.

We use the results of this analysis as reference abundances for the experiments done in Section 3.5.



**Fig. 3.** Comparison of F-scores (combination of precision and sensitivity) obtained from different mixes for different combinations of methods, data type and databases. Red dots indicate a  $p$ -value below 0.05. Combinations of methods, data type and databases are shown on the horizontal axis, Taxonomic levels are shown on the vertical axis. RS – RefSeq database, GG – GreenGenes database, SILVA – SILVA database. (For interpretation of the references to colour in this figure legend, the reader is referred to the web version of this article.)

### 3.3. Analysis of bacterial mixes using Centrifuge and MG-RAST

The mixes were sequenced on the Illumina MiSeq using WGS (datasets EQ\_WGS, EXP16\_WGS, EXP12\_WGS and EXP56\_WGS) and 16S for V3-V4 region (datasets EQ\_16S, EXP16\_16S, EXP12\_16S and EXP56\_16S). Information about read counts and QC statistics for each obtained dataset can be found in Supplementary Table S1.

WGS and 16S datasets obtained from our four metagenomic mixes were analysed with Centrifuge using the RefSeq complete bacterial

genomes database. We performed additional analysis for 16S datasets using Centrifuge with the GreenGenes and SILVA reference databases.

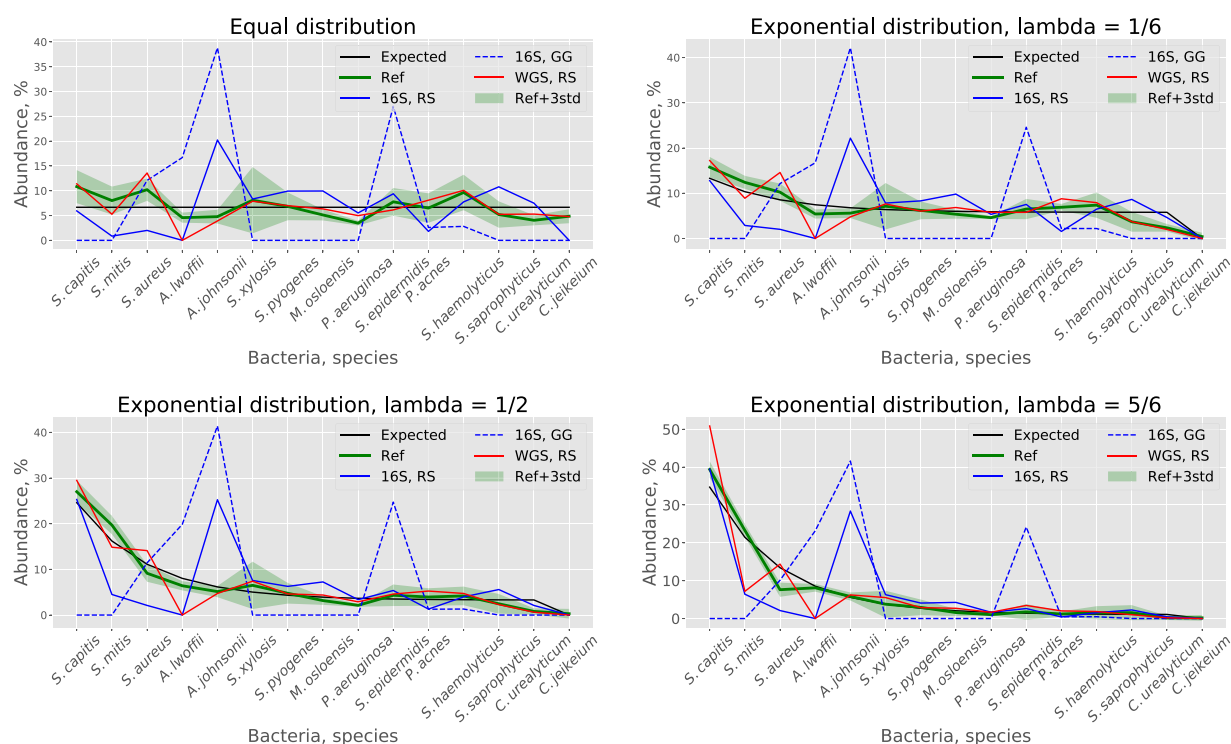
All eight datasets four WGS and four 16S were submitted to the MG-RAST Metagenomics analysis server under project number 85582. RefSeq and GreenGenes databases provide taxonomic annotation down to the species level, while SILVA database as well as the databases used by MG-RAST are restricted to the genus level. Since the NCBI taxonomy and the taxonomy used by MG-RAST were different at the order level for our set of bacteria, we excluded annotation at the order level from further analysis.

### 3.4. Profiling accuracy without considering relative abundances

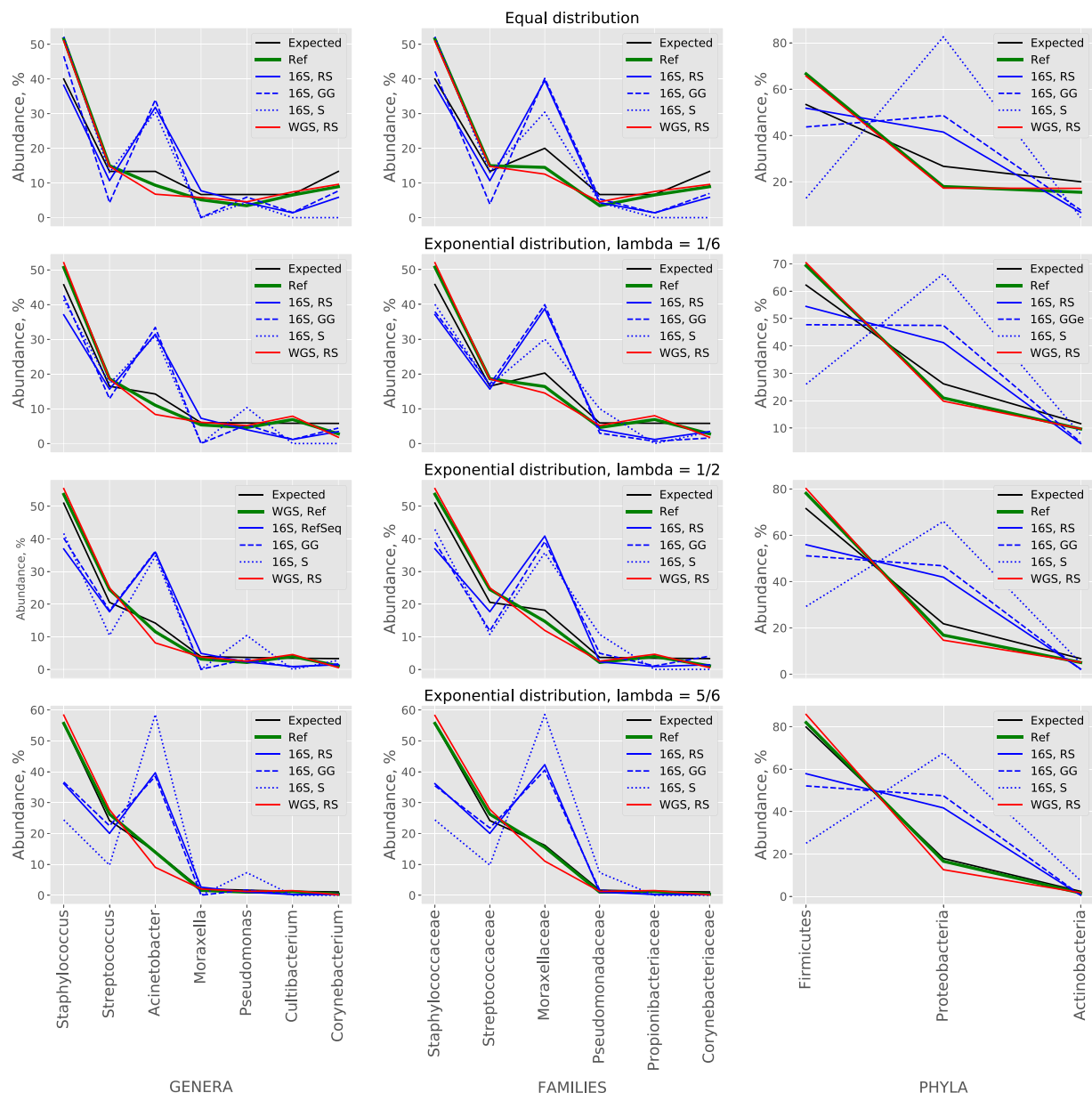
Because the content of the metagenomic mixes is known, we can verify how many of the reported taxa on each taxonomic rank are correct (true positive counts), how many are incorrect (false positive counts) and how many are missed (false negative counts).

Using these counts, both precision and sensitivity can be calculated. A perfect prediction is made if both precision and sensitivity equal one. As can be seen in Fig. 2, both precision and sensitivity tend to increase in all cases with increasing taxonomic rank. For all 16S datasets analysed with Centrifuge, we observe that precision never reaches its maximum value, while for WGS datasets analysed with Centrifuge precision reaches its maximum already at the genus level. Interestingly, for 16S datasets analysed with MG-RAST, precision reaches its maximum at the genus level, but the sensitivity does not increase any further. For WGS datasets analysed with MG-RAST, sensitivity reaches its maximum already at the family level.

The accuracy of the classifications can be expressed using the F-score, a measure of prediction accuracy that considers both precision and sensitivity. We tested whether the F-scores differed significantly for each pair-wise comparison using the Mann-Whitney  $U$  test and the Benjamini-Hochberg procedure for FDR control. The full table of  $p$ -values can be found in Supplementary Table S2, a summary of the



**Fig. 4.** Comparison of relative abundances reported by Centrifuge (using two different reference databases) for WGS and 16S with relative abundances obtained from regression analysis. Results for each mix are shown separately with the species' names on the horizontal axis and the relative abundances on the vertical axis. Ref – reference abundances, RS – RefSeq database, GG – Greengenes database. Please note that data points are connected only to visualize the various types of distributions.



**Fig. 5.** Comparison of relative abundances reported by Centrifuge (using three different reference databases) for WGS and 16S datasets on genera, orders and phyla levels with relative reference abundances. In the above grid of figures each row indicates the mix and each column indicates the taxonomic level. In each figure, the taxa are shown on the horizontal axis and the relative abundances are shown on the vertical axis. Ref – reference abundances, RS – RefSeq database, GG – Greengenes database, S - SILVA database.

results is shown in Fig. 3. In most cases, the F-scores differ significantly when comparing WGS to 16S.

### 3.5. Abundance assignment accuracy

Both Centrifuge and MG-RAST provide read counts for each reported taxon. We considered only reads that were assigned to the expected taxa and compared their relative abundances to the reference abundances.

Only Centrifuge, when using either the RefSeq or GreenGenes database, reported the taxonomic assignment down to the species level. In Fig. 4, each metagenomic mix is shown as a separate graph with species listed on the horizontal axes and their relative abundances shown on the vertical axes. The black line represents the intended distribution of species abundances. The dark green line shows the mean reference abundances with the light green area representing  $\pm 3$  standard

deviation around those means. The blue and red lines show the relative abundances obtained for 16S and WGS datasets respectively, with the solid blue line for the 16S analysis done using the RefSeq database and the dashed blue line using the GreenGenes database. As can be seen in Fig. 4, the analysis of 16S data results in a considerable overestimation of abundance of *A. johnsonii*. Centrifuge failed to identify *A. lwoffii*, since there is no complete genome of that bacterium in the RefSeq database and it did not report any significant presence of *C. jeikeium* in the exponentially distributed metagenomic mixes. Analysis of the 16S datasets using the GreenGenes database reported overestimated values for *S. epidermidis* and *A. johnsonii* and did not report the presence of nine out of fifteen bacteria because of their absence in the GreenGenes database.

We repeated the same analysis on three higher taxonomic ranks: genera, families and phyla. For all these three taxonomic levels we analysed the results of Centrifuge (Fig. 5) and MG-RAST (Fig. 6). As can

be seen in Fig. 4, the Centrifuge analysis of 16S datasets using different reference databases provided a similar biased output, mostly due to an overestimation of the abundance of the *Acinetobacter* genus, *Moraxelaceae* family and *Proteobacteria* phylum. The dissimilarity with the reference abundances is especially pronounced at the phylum level. Results obtained for the WGS datasets with Centrifuge were concordant with the reference abundances with slight deviation for *Acinetobacter* genus, *Moraxelaceae* family and *Proteobacteria* phylum (Fig. 5). It is interesting to note, that these taxa were also the major reason for disagreement between results obtained by Centrifuge for 16S datasets and reference abundances.

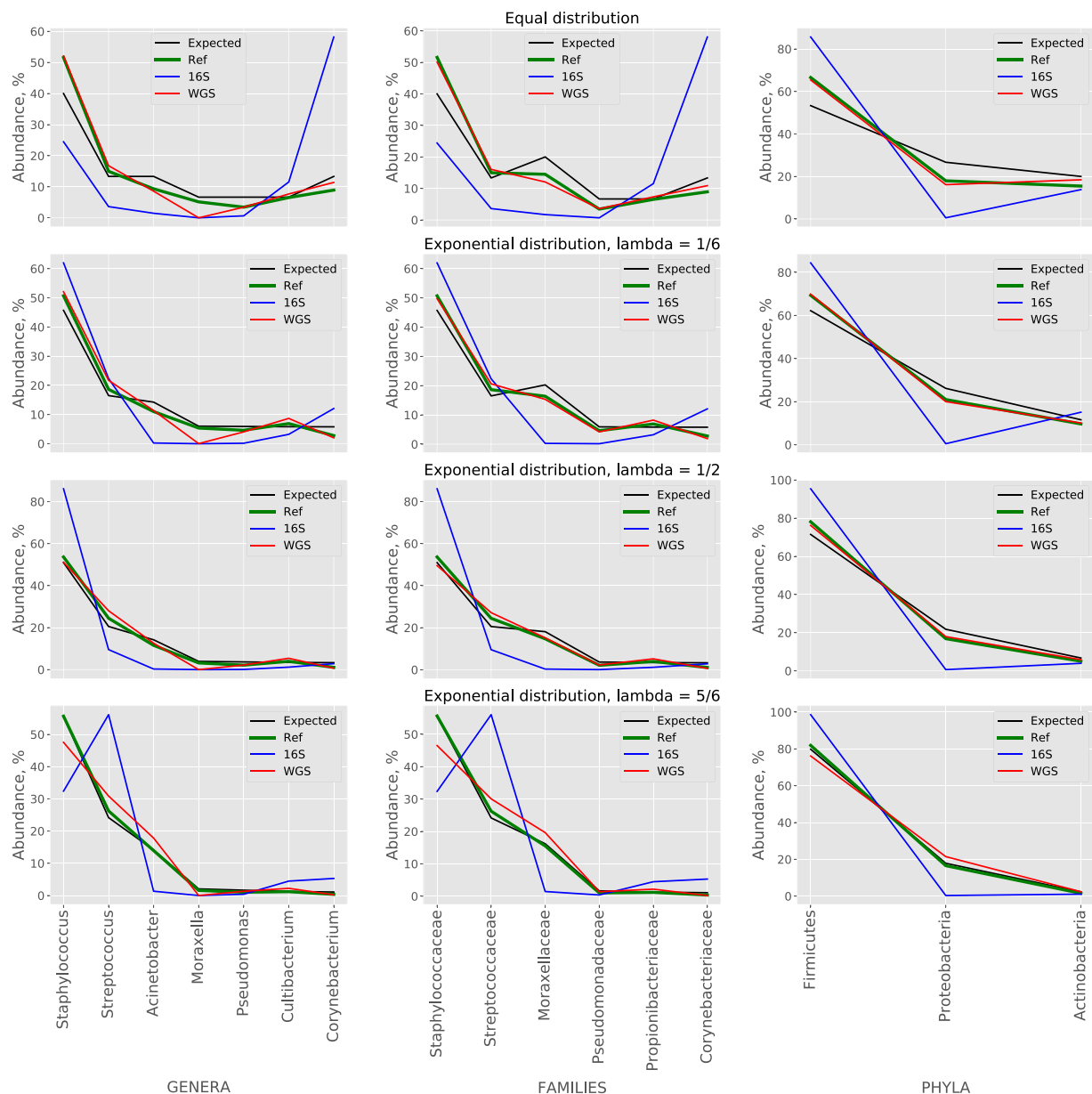
The results obtained for different 16S datasets by MG-RAST were not consistent (as is the case for Centrifuge) up to the phylum level. As can be seen in Fig. 5, analysis of 16S datasets with MG-RAST reported many disagreements with reference abundances. The reasons of those disagreements are dataset- and taxonomy rank-specific. Results reported by MG-RAST became more or less consistent only at the phylum

level, where they followed the same trend: overestimating the abundance of Firmicutes relative to that of Proteobacteria.

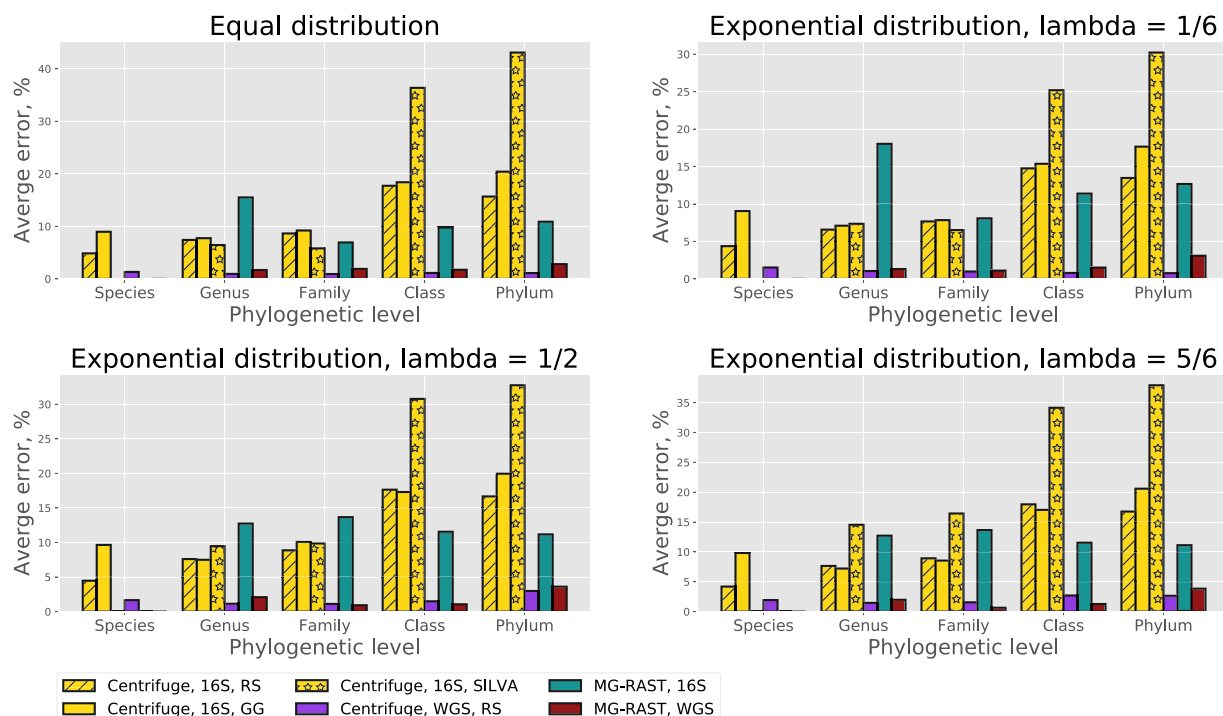
Abundances obtained after analysis with MG-RAST of WGS datasets were also following the reference results closely. There were, however, slight deviations from the reference abundances. These deviations were, like the results for 16S datasets, specific to taxonomy-rank and dataset.

In order to quantify the dissimilarity among the abundances provided by the different methods, datasets, reference databases and the results of regression analysis we calculated the absolute differences in abundances for each particular dataset and taxonomic rank. The averages of these values (from here on called the error rate) are reported in Fig. 7.

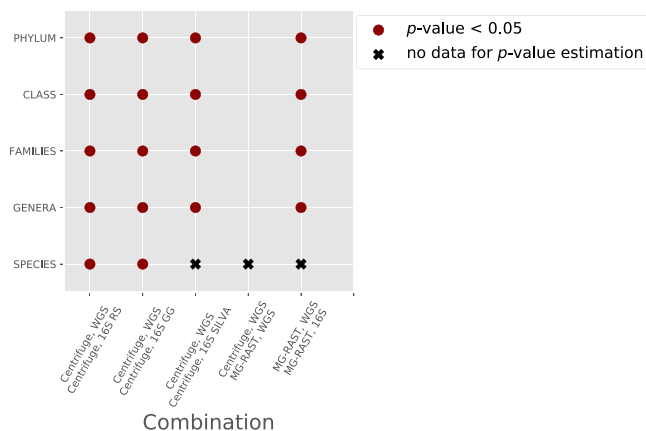
For the analyses of 16S datasets it is interesting to note that for Centrifuge the average error rate grew with the increase of the taxonomic rank in general. This was not the case for the error rate obtained for the 16S datasets using MG-RAST. We tested whether the average errors differed significantly for each pair-wise comparison using the



**Fig. 6.** Comparison of relative abundances reported by MG-RAST for WGS and 16S datasets on genera, orders and phyla levels with relative reference abundances. In the above grid of figures each row indicates the mix and each column indicates the taxonomic level. In each figure, the taxa are shown on the horizontal axis and the relative abundances are shown on the vertical axis. Ref – reference abundances.



**Fig. 7.** Average error between the reference abundances and results obtained for different combinations of datasets, reference databases and tools. Results for each metagenomics mix are shown separately. Bar colors represent the tool, data type and reference database combination used for the analysis. The error rate (shown on the vertical axes) was calculated for each taxonomic rank (shown on the horizontal axes) separately. RS – RefSeq database, GG – GreenGenes database, SILVA – SILVA database.



**Fig. 8.** Comparison of average errors obtained from different mixes for different combinations of methods, data type and databases. Red dots indicate a  $p$ -value below 0.05. Combinations of methods, data type and databases are shown on the horizontal axis, Taxonomic levels are shown on the vertical axis. RS – RefSeq database, GG – GreenGenes database, SILVA – SILVA database. (For interpretation of the references to colour in this figure legend, the reader is referred to the web version of this article.)

Mann-Whitney  $U$  test and the Benjamini-Hochberg procedure for FDR control. The full table of  $p$ -values can be found in Supplementary Table S3, a summary of the results is shown in Fig. 8. This analysis demonstrates that for all taxonomic levels the error rates in the abundance estimations provided by the analysis of 16S datasets (regardless of the method or reference database) are significantly different (higher) compared to the abundances reported for WGS datasets. We did not observe any significant difference in average error rate between WGS datasets analysed with Centrifuge and MG-RAST.

We compared the error rates reported by Centrifuge when using the three different 16S reference databases. Error rates observed in the

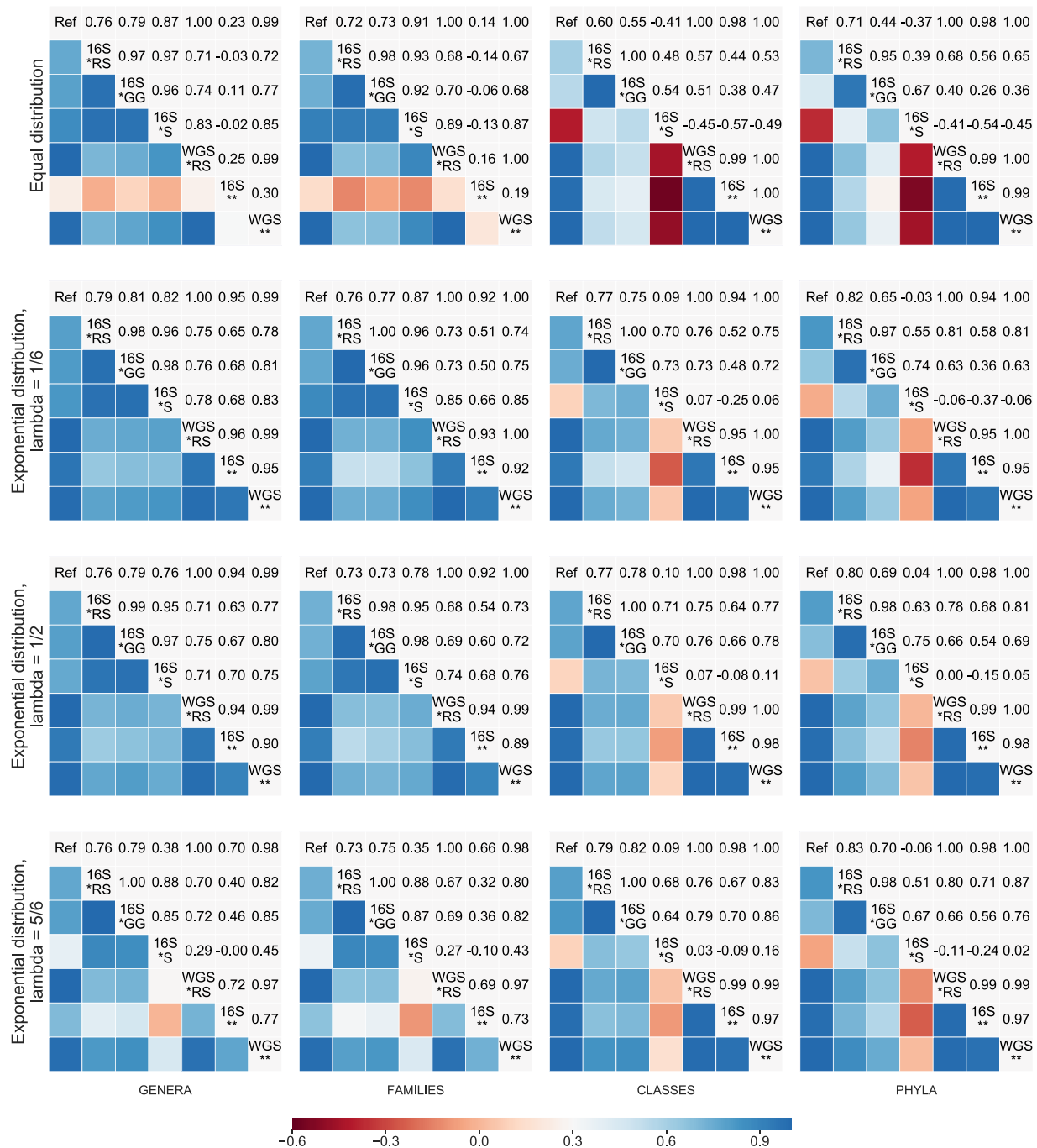
analysis with RefSeq and GreenGenes databases were similar. Running the Centrifuge analysis using the SILVA database reported a much higher error rate. That might be a direct consequence of taxonomic annotation done using the SILVA database where a smaller proportion of reads was assigned to the expected taxa in comparison to other reference databases (see Section 3.4).

We also evaluated the similarity among the abundances obtained by employing distinct methods and databases using a correlation analysis. In Fig. 9 the results of these comparisons are presented as a series of heatmaps. As can be seen from Fig. 9, abundances obtained by the analysis of WGS data (Centrifuge and MG-RAST) for all datasets at all taxonomic levels positively correlate with reference abundances. Correlation of 16S analysis obtained using Centrifuge with the reference abundances becomes worse at higher taxonomic levels, which is the opposite for the 16S data results obtained using MG-RAST. The 16S data analyses obtained for Centrifuge and MG-RAST do not demonstrate positive correlation with each other.

#### 4. Conclusions

In this study we created a series of bacterial mixes with known content in order to investigate which type of metagenomics data and reads assignment strategy yields better taxonomic classification. For each mix we generated WGS and 16S sequencing datasets and analysed them using Centrifuge with RefSeq, GreenGenes and SILVA reference databases and the MG-RAST metagenomics analysis server with M5nr and M5nr reference databases. We compared the results of all analysis done with Centrifuge and MG-RAST to the reference abundance profiles obtained from a  $k$ -mer-based regression analysis.

The results from both Centrifuge and MG-RAST show that WGS datasets provide much more accurate results in comparison to 16S-based methods. The analysis of WGS data displayed better coverage of all taxa expected to be present in the mixes on all phylogenetic levels, reaching the maximum accuracy already at the genus level for Centrifuge and at the family level for MG-RAST. On the other hand,



**Fig. 9.** Correlation between the abundance profiles obtained for different combinations of method, dataset and reference database on distinct taxonomic levels. In the above grid of figures each row indicates the profiles obtained for different combinations of method, dataset and reference database on distinct taxonomic levels. In the above grid of figures each row indicates the profiles obtained for different combinations of method, dataset and reference database on distinct taxonomic levels. In the above grid of figures each row indicates the profiles obtained for different combinations of method, dataset and reference database on distinct taxonomic levels. The combination of analysis type, dataset and reference database are shown on the main diagonal of the heatmap, with the lower triangle representing the correlation shown in colors and the upper triangle demonstrating the same data in the numeric representation. Ref – reference abundances, \* – Centrifuge, \*\* – MG-RAST, RS – RefSeq database, GG – GreenGenes database, S – SILVA database.

results obtained for 16S-based data were often missing several taxa and/or had very high false-positive rates. Centrifuge analyses based on the 16S datasets were suffering from low precision, while MG-RAST analysis of the 16S datasets had low sensitivity. Abundance profiles obtained from WGS demonstrated much less disagreement with the expected abundances in comparison to the abundance profiles based on 16S data. This was shown using two different measurements: the average (per taxonomy rank) absolute difference between abundance profiles and by a correlation analysis.

For 16S datasets analysed with Centrifuge, the deviation from the reference abundances, introduced at the species/genus levels,

propagated further up the taxonomy which led to a greater difference with the expected outcome on the higher taxonomic ranks as well. In contrast, the analysis of 16S datasets performed by the MG-RAST pipeline demonstrated greater differences with the reference abundances on the lower taxonomic ranks in comparison with the higher ones. Our correlation analysis shows that the agreement between the MG-RAST results of 16S datasets and reference abundances was growing with increasing taxonomic level.

Both tailor-made 16S databases (GreenGenes and SILVA) did not perform better than the RefSeq database when analysing 16S datasets using Centrifuge. The Centrifuge results using RefSeq and GreenGenes

databases were correlated with a correlation coefficient higher than 0.95 for all 16S datasets on each taxonomic rank starting with genus.

We conclude that WGS data is preferable for the study of metagenomic data, especially when the correct inhabitant abundances are required. We could not determine which of the explored methods for the taxonomic assignment of the WGS data provides a more accurate outcome. Centrifuge, however, has minor advantages in comparison to MG-RAST, such as a faster, deeper and slightly better reads classification, the possibility of local installation and use of custom databases and a more flexible tuning of the tools' settings. Among the investigated techniques for 16S metagenomic data analysis, MG-RAST demonstrated slightly better results in both reads assignment and abundance estimation, albeit only at higher taxonomic ranks.

As previously quoted, "the capacity of WGS data of microbiomes to aid in forensic investigations by connecting objects and environments to individuals has been poorly investigated". In light of this, our results are especially important, as they demonstrate the inefficiency of routine 16S data to produce the accurate taxonomical profiling.

The synthetic metagenomes created in our study is restricted to DNA of bacteria that inhabit skin surface – a logical target for forensics analysis. However, human skin is also the environment with one of the most within- and between-individual diverse microbiota on the human body. The benchmark we created is rather small and simple as the diversity of microbial species living on the human skin surface is much larger than only 15 species [95]. The significant inaccuracy of the results obtained for 16S data in comparison with those for WGS data on a small and simple set of benchmarks can possibly question the accuracy of the previous 16S-based forensic studies, at least those done on skin-associated microbial communities.

## Data accessibility

Reads obtained after the individual sequencing of each selected bacterial species and used for the genome assembly were upload to the Sequence Read Archive under the study SRP159200.

Sequencing reads of the metagenomic mixes as well as the results of the analysis performed by MG-RAST can be downloaded from the MG-RAST server project number 85582.

The summary of results obtained using Centrifuge as well as the supplementary materials for this research can be found on Figshare: <https://doi.org/10.6084/m9.figshare.c.4217672>.

## Funding

This research is financed by a grant number 727.011.002 of the Netherlands Organisation for Scientific Research (NWO).

## Acknowledgments

We would like to thank Guy Allard for the support with the assembly of bacterial genomes and Louk Rademaker for the feedback on this manuscript.

## Appendix A. Supplementary data

Supplementary material related to this article can be found, in the online version, at doi:<https://doi.org/10.1016/j.fsigen.2020.102257>.

## References

- [1] C. Rinke, P. Schwientek, A. Sczyrba, N.N. Ivanova, et al., Insights into the phylogeny and coding potential of microbial dark matter, *Nature* 499 (7459) (2013) 431–437, <https://doi.org/10.1038/nature12352>.
- [2] C.E. Mason, S. Tighe, Focus on metagenomics, *J. Biomol. Tech.* 28 (1) (2017) 1, <https://doi.org/10.17171/jbt.17-2801-010>.
- [3] S. Kumar, K.K. Krishnani, B. Bhushan, M.P. Brahmane, Metagenomics: retrospect and prospects in high throughput age, *Biotechnol. Res. Int.* 2015 (2015) 121735, <https://doi.org/10.1155/2015/121735>.
- [4] J.T. Staley, A. Konopka, Measurement of in situ activities of non-photosynthetic microorganisms in aquatic and terrestrial habitats, *Annu. Rev. Microbiol.* 39 (1) (1985) 321–346.
- [5] D.B. Roszak, R.R. Colwell, Survival strategies of bacteria in the natural environment, *Microbiol. Rev.* 51 (3) (1987) 365–379.
- [6] E.J. Stewart, Growing unculturable bacteria, *J. Bacteriol.* 194 (16) (2012) 4151–4160, <https://doi.org/10.1128/JB.00345-12>.
- [7] K.I. Mohr, Diversity of myxobacteria – we only see the tip of the iceberg, *Microorganisms* 6 (3) (2018), <https://doi.org/10.3390/microorganisms6030084>.
- [8] M. Hamady, C. Fraser-Liggett, R. Knight, et al., The human microbiome project: Exploring the microbial part of ourselves in a changing world, *Nature* 449 (7164) (2007) 804–810, <https://doi.org/10.1038/nature06244>.
- [9] W.-L. Wang, S.-Y. Xu, Z.-G. Ren, L. Tao, J.-W. Jiang, et al., Application of metagenomics in the human gut microbiome, *World J. Gastroenterol.* 21 (3) (2015) 803–814, <https://doi.org/10.3748/wjg.v21.i3.803>.
- [10] S. Al Khodor, B. Reichert, I.F. Shatat, The microbiome and blood pressure: can microbes regulate our blood pressure? *Front. Pediatr.* 5 (2017) 138, <https://doi.org/10.3389/fped.2017.00138>.
- [11] R. Kolde, E.A. Franzosa, G. Rahnnavard, A.B. Hall, H. Vlamakis, et al., Host genetic variation and its microbiome interactions within the human microbiome project, *Genome Med.* 10 (1) (2018) 6, <https://doi.org/10.1186/s13073-018-0515-8>.
- [12] M. Hattori, T.D. Taylor, The human intestinal microbiome: a new frontier of human biology, *DNA Res.* 16 (1) (2009) 1–12, <https://doi.org/10.1093/dnares/dsn033>.
- [13] A. Kouzuma, S. Ishii, K. Watanabe, Metagenomic insights into the ecology and physiology of microbes in bioelectrochemical systems, *Bioresour. Technol.* 255 (2018) 302–307, <https://doi.org/10.1016/j.biortech.2018.01.125>.
- [14] F.H. Coutinho, G.B. Gregoracci, J.M. Walter, C.C. Thompson, F.L. Thompson, Metagenomics sheds light on the ecology of marine microbes and their viruses, *Trends Microbiol.* 26 (11) (2018) 955–965, <https://doi.org/10.1016/j.tim.2018.05.015>.
- [15] R.J. Boissy, D.J. Romberger, W.A. Roughead, L. Weissenburger-Moser, J.A. Poole, et al., Shotgun pyrosequencing metagenomic analyses of dusts from swine confinement and grain facilities, *PLoS One* 9 (4) (2014) e95578, <https://doi.org/10.1371/journal.pone.0095578>.
- [16] B. Carbonetto, N. Rascovan, R.A. Ivarez, A. Mentaberry, M.P. Vazquez, Structure, composition and metagenomic profile of soil microbiomes associated to agricultural land use and tillage systems in Argentine pampas, *PLoS One* 9 (6) (2014) e99949, <https://doi.org/10.1371/journal.pone.0099949>.
- [17] J.Q. Su, B. Wei, C.Y. Xu, M. Qiao, Y.G. Zhu, Functional metagenomic characterization of antibiotic resistance genes in agricultural soils from China, *Environ. Int.* 65 (2014) 9–15, <https://doi.org/10.1016/j.envint.2013.12.010>.
- [18] S.J. Finley, M.E. Benbow, G.T. Javan, Microbial communities associated with human decomposition and their potential use as postmortem clocks, *Int. J. Legal Med.* 129 (3) (2015) 623–632, <https://doi.org/10.1007/s00414-014-1059-0>.
- [19] A. Fornaciari, Environmental microbial forensics and archaeology of past pandemics, *Micro. Biol. Spectrum* 5 (1) (2017), <https://doi.org/10.1128/microbiolspec.EMF-0011-2016>.
- [20] N.A. Bokulich, Z.T. Lewis, K. Boundy-Mills, D.A. Mills, A new perspective on microbial landscapes within food production, *Curr. Opin. Biotechnol.* 37 (2016) 182–189, <https://doi.org/10.1016/j.copbio.2015.12.008>.
- [21] Q. Peng, X. Wang, M. Shang, J. Huang, G. Guan, et al., Isolation of a novel alkaline-stable lipase from a metagenomic library and its specific application for milkfat flavor production, *Microb. Cell Fact.* 13 (1) (2014) 1, <https://doi.org/10.1186/1475-2859-13-1>.
- [22] M. Trindade, L.J. van Zyl, J. Navarro-Fernandez, A. Abd Elrazak, Targeted metagenomics as a tool to tap into marine natural product diversity for the discovery and production of drug candidates, *Front. Microbiol.* 6 (2015) 890, <https://doi.org/10.3389/fmicb.2015.00890>.
- [23] M. Drancourt, C. Bollet, A. Carliz, Martelin, J.-P. Gayral, D. Raoult, 16S ribosomal DNA sequence analysis of a large collection of environmental and clinical unidentified bacterial isolates, *J. Clin. Microbiol.* 38 (10) (2000) 3623–3630.
- [24] G. Muyzer, E.C. De Waal, A.G. Uitterlinden, Profiling of complex microbial populations by denaturing gradient gel electrophoresis analysis of polymerase chain reaction amplified genes coding for 16S rRNA, *Appl. Environ. Microbiol.* 59 (3) (1993) 695–700.
- [25] C.R. Woese, O. Kandler, M.L. Wheelis, Towards a natural system of organisms: proposal for the domains archaea, bacteria, and eucarya, *Proc. Natl. Acad. Sci.* 87 (12) (1990) 4576–4579, <https://doi.org/10.1073/pnas.87.12.4576>.
- [26] C.R. Woese, G.E. Fox, Phylogenetic structure of the prokaryotic domain: the primary kingdoms, *Proc. Natl. Acad. Sci.* 74 (11) (1977) 5088–5090, <https://doi.org/10.1073/pnas.74.11.5088>.
- [27] M. Balvociute, D.H. Huson, SILVA, RDP, GreenGenes, NCBI and OTT – how do these taxonomies compare? *BMC Genomics* 18 (2) (2017) 114, <https://doi.org/10.1186/s12864-017-3501-4>.
- [28] G.J. Olsen, R. Overbeek, N. Larsen, T.L. Marsh, M.J. McCaughey, et al., The ribosomal database project, *Nucleic Acids Res. (Suppl.)* 20 (1992) 2199–2200, <https://doi.org/10.1093/nar/22.17.3485>.
- [29] J.R. Cole, Q. Wang, J.A. Fish, B. Chai, D.M. McFarrell, et al., Ribosomal database project: data and tools for high throughput rRNA analysis, *Nucleic Acids Res.* 42 (D1) (2013) D633–D642, <https://doi.org/10.1093/nar/gkt1244>.
- [30] D. McDonald, M.N. Price, J. Goodrich, E.P. Nawrocki, T.Z. DeSantis, et al., An improved GreenGenes taxonomy with explicit ranks for ecological and evolutionary analyses of bacteria and archaea, *ISME J.* 6 (3) (2012) 610, <https://doi.org/10.1038/ismej.2011.139>.
- [31] P. Yilmaz, L.W. Parfrey, P. Yarza, J. Gerken, E. Pruesse, et al., The SILVA and "all-

- species living tree project (ltp)" taxonomic frameworks, *Nucleic Acids Res.* 42 (D1) (2013) D643–D648, <https://doi.org/10.1093/nar/gkt1209>.
- [32] J.M. Janda, S.L. Abbott, 16S rRNA gene sequencing for bacterial identification in the diagnostic laboratory: pluses, perils, and pitfalls, *J. Clin. Microbiol.* 45 (9) (2007) 2761–2764, <https://doi.org/10.1128/JCM.01228-07>.
- [33] J.-H. Ahn, B.-Y. Kim, J. Song, H.-Y. Weon, Effects of PCR cycle number and DNA polymerase type on the 16S rRNA gene pyrosequencing analysis of bacterial communities, *J. Microbiol.* 50 (6) (2012) 1071–1074, <https://doi.org/10.1007/s12275-012-2642-z>.
- [34] J.P. Brooks, D.J. Edwards, M.D. Harwich, M.C. Rivera, J.M. Fettweis, et al., The truth about metagenomics: quantifying and counteracting bias in 16S rRNA studies, *BMC Microbiol.* 15 (1) (2015) 66, <https://doi.org/10.1186/s12866-015-0351-6>.
- [35] A.J. Pinto, L. Raskin, PCR biases distort bacterial and archaeal community structure in pyrosequencing datasets, *PLoS One* 7 (8) (2012) e43093, <https://doi.org/10.1371/journal.pone.0043093>.
- [36] S. Hong, J. Bunge, C. Leslin, S. Jeon, S.S. Epstein, Polymerase chain reaction primers miss half of rRNA microbial diversity, *ISME J.* 3 (12) (2009) 1365, <https://doi.org/10.1038/ismej.2009.89>.
- [37] P.H.A. Timmers, H.C.A. Widjaja-Greefkes, C.M. Plugge, A.J.M. Stams, Evaluation and optimization of PCR primers for selective and quantitative detection of marine ANME subclusters involved in sulfate-dependent anaerobic methane oxidation, *Appl. Microbiol. Biotechnol.* 101 (14) (2017) 5847–5859, <https://doi.org/10.1007/s00253-017-8338-x>.
- [38] S. Ceuppens, D. De Coninck, N. Botteldoorn, F. Van Nieuwerburgh, M. Uyttendaele, Microbial community profiling of fresh basil and pitfalls in taxonomic assignment of enterobacterial pathogenic species based upon 16S rRNA amplicon sequencing, *Int. J. Food Microbiol.* 257 (2017) 148–156, <https://doi.org/10.1016/j.ijfoodmicro.2017.06.016>.
- [39] F.E. Dewhirst, Z. Shen, M.S. Scimeca, L.N. Stokes, T. Boumenna, et al., Discordant 16S and 23S rRNA gene phylogenies for the genus *helicobacter*: implications for phylogenetic inference and systematics, *J. Bacteriol.* 187 (17) (2005) 6106–6118, <https://doi.org/10.1128/JB.187.17.6106-6118.2005>.
- [40] S.W. Kembel, M. Wu, J.A. Eisen, J.L. Green, Incorporating 16S gene copy number information improves estimates of microbial diversity and abundance, *PLoS Comput. Biol.* 8 (10) (2012) e1002743, <https://doi.org/10.1371/journal.pcbi.1002743>.
- [41] E.N. Hanssen, K.H. Liland, P. Gill, L. Snipen, Optimizing body fluid recognition from microbial taxonomic profiles, *Forensic Sci. Int. Genet.* 37 (2018) 13–20, <https://doi.org/10.1016/j.fsigen.2018.07.012>.
- [42] N. Fierer, C.L. Lauber, N. Zhou, D. McDonald, E.K. Costello, R. Knight, Forensic identification using skin bacterial communities, *Proc. Natl. Acad. Sci.* 107 (14) (2010) 6477–6481, <https://doi.org/10.1073/pnas.1000162107>.
- [43] S. Lax, J.T. Hampton-Marcell, S.M. Gibbons, G.B. Colares, D. Smith, et al., Forensic analysis of the microbiome of phones and shoes, *Microbiome* 3 (1) (2015) 21, <https://doi.org/10.1186/s40168-015-0082-9>.
- [44] G.J. Olsen, C.R. Woese, Ribosomal RNA: a key to phylogeny, *FASEB J.* 7 (1) (1993) 113–123, <https://doi.org/10.1096/asebj.7.1.8422957>.
- [45] P. Yilmaz, R. Kottmann, E. Pruesse, C. Quast, F.O. Glöckner, Analysis of 23S rRNA genes in metagenomes – a case study from the Global Ocean Sampling Expedition, *Syst. Appl. Microbiol.* 34 (6) (2011) 462–469, <https://doi.org/10.1016/j.syapm.2011.04.005>.
- [46] L. Yang, Z. Tan, D. Wang, L. Xue, M.-X. Guan, T. Huang, R. Li, Species identification through mitochondrial rRNA genetic analysis, *Sci. Rep.* 4 (2014) 4089, <https://doi.org/10.1038/srep04089> EP.
- [47] A.E. Budding, M. Martine Hoogewerf, C.M.J.E. Vandenbroucke-Grauls, P.H.M. Savelkoul, Automated broad-range molecular detection of bacteria in clinical samples, *J. Clin. Microbiol.* 54 (4) (2016) 934–943, <https://doi.org/10.1128/JCM.02886-15>.
- [48] F.C. Quak, T.V. Duijn, J. Hoogenboom, A.D. Kloosterman, I. Kuiper, Human-associated microbial populations as evidence in forensic casework, *Forensic Sci. Int. Genet.* 36 (2018) 176–185, <https://doi.org/10.1016/j.fsigen.2018.06.020>.
- [49] S.E. Schmedes, A.E. Woerner, N.M. Novroski, F.R. Wendt, J.L. King, et al., Targeted sequencing of clade-specific markers from skin microbiomes for forensic human identification, *Forensic Sci. Int. Genet.* 32 (2018) 50–61, <https://doi.org/10.1016/j.fsigen.2017.10.004>.
- [50] A.E. Woerner, N.M. Novroski, F.R. Wendt, A. Ambers, R. Wiley, et al., Forensic human identification with targeted microbiome markers using nearest neighbor classification, *Forensic Sci. Int. Genet.* 38 (2019) 130–139, <https://doi.org/10.1016/j.fsigen.2018.10.003>.
- [51] R. Ranjan, A. Rani, A. Metwally, H.S. McGee, D.L. Perkins, Analysis of the microbiome: advantages of whole genome shotgun versus 16S amplicon sequencing, *Biochem. Biophys. Res. Commun.* 469 (4) (2016) 967–977, <https://doi.org/10.1016/j.bbrc.2015.12.083>.
- [52] J. Jovel, J. Patterson, W. Wang, N. Hotte, S. O'Keefe, et al., Characterization of the gut microbiome using 16S or shotgun metagenomics, *Front. Microbiol.* 7 (2016) 459, <https://doi.org/10.3389/fmicb.2016.00459>.
- [53] N. Shah, H. Tang, T.G. Doak, Y. Ye, Comparing bacterial communities inferred from 16S rRNA gene sequencing and shotgun metagenomics, *Bioinformatics* 2010 (2011) 165–176, [https://doi.org/10.1142/9789814335058\\_0018](https://doi.org/10.1142/9789814335058_0018).
- [54] B. Steven, L.V. Gallegos-Graves, S.R. Starkenburg, P.S. Chain, C.R. Kuske, Targeted and shotgun metagenomic approaches provide different descriptions of dryland soil microbial communities in a manipulated field study, *Environ. Microbiol. Rep.* 4 (2) (2012) 248–256, <https://doi.org/10.1111/j.1758-2229.2012.00328.x>.
- [55] A. Almeida, A.L. Mitchell, A. Tarkowska, R.D. Finn, Benchmarking taxonomic assignments based on 16S rRNA gene profiling of the microbiota from commonly sampled environments, *Gigascience* 7 (5) (2018), <https://doi.org/10.1093/gigascience/giy054> giy054.
- [56] T. Sijen, Molecular approaches for forensic cell type identification: on mRNA, miRNA, DNA methylation and microbial markers, *Forensic Sci. Int. Genet.* 18 (2015) 21–32, <https://doi.org/10.1016/j.fsigen.2014.11.015>.
- [57] T.H. Clarke, A. Gomez, H. Singh, K.E. Nelson, L.M. Brinkac, Integrating the microbiome as a resource in the forensics toolkit, *Forensic Sci. Int. Genet.* 30 (2017) 141–147, <https://doi.org/10.1016/j.fsigen.2017.06.008>.
- [58] R. D'Amore, U.Z. Ijaz, M. Schirmer, J.R. Kenny, R. Gregory, et al., A comprehensive benchmarking study of protocols and sequencing platforms for 16S rRNA community profiling, *BMC Genomics* 17 (1) (2016), <https://doi.org/10.1186/s12864-015-2194-9>.
- [59] A. Szczyrba, P. Hofmann, P. Belmann, et al., Critical Assessment of Metagenome Interpretation—a benchmark of metagenomics software, *Nat. Methods* 14 (11) (2017) 1063–1071, <https://doi.org/10.1038/nmeth.4458>.
- [60] M.A. Peabody, T. Van Rossum, R. Lo, F.S. Brinkman, Evaluation of shotgun metagenomics sequence classification methods using in silico and in vitro simulated communities, *BMC Bioinformatics* 16 (2015) 363, <https://doi.org/10.1186/s12859-015-0788-5>.
- [61] K. Mavromatis, N. Ivanova, K. Barry, H. Shapiro, E. Goltsman, et al., Use of simulated data sets to evaluate the fidelity of metagenomic processing methods, *Nat. Methods* 4 (6) (2007) 495–500, <https://doi.org/10.1038/nmeth1043>.
- [62] A.B.R. McIntyre, R. Ounit, E. Afshinnikoo, et al., Comprehensive benchmarking and ensemble approaches for metagenomic classifiers, *Genome Biol.* 18 (1) (2017) 182, <https://doi.org/10.1186/s13059-017-1299-7>.
- [63] A.M. Walsh, F. Crispie, O. O'Sullivan, L. Finnegan, M.J. Claesson, et al., Species classifier choice is a key consideration when analysing low-complexity food microbiome data, *Microbiome* 6 (1) (2018) 50, <https://doi.org/10.1186/s40168-018-0437-0>.
- [64] N. Segata, L. Waldron, A. Ballarín, V. Narasimhan, O. Jousson, C. Huttenhower, Metagenomic microbial community profiling using unique clade-specific marker genes, *Nat. Methods* 9 (8) (2012) 811–814, <https://doi.org/10.1038/nmeth.2066>.
- [65] V.C. Piro, M. Matuszkowski, B.Y. Renard, MetaMeta: integrating metagenome analysis tools to improve taxonomic profiling, *Microbiome* 5 (1) (2017) 101, <https://doi.org/10.1186/s40168-017-0318-y>.
- [66] A. Escobar-Zepeda, E.E. Godoy-Lozano, L. Raggi, L. Segovia, E. Merino, et al., Analysis of sequencing strategies and tools for taxonomic annotation: defining standards for progressive metagenomics, *Sci. Rep.* 8 (1) (2018) 12034, <https://doi.org/10.1038/s41598-018-30515-5>.
- [67] S. Lindgreen, K.L. Adair, P.P. Gardner, An evaluation of the accuracy and speed of metagenome analysis tools, *Sci. Rep.* 6 (1) (2016) 19233, <https://doi.org/10.1038/srep19233>.
- [68] D. Kim, L. Song, F.P. Breitwieser, S.L. Salzberg, Centrifuge: rapid and sensitive classification of metagenomic sequences, *Genome Res.* 26 (2016) 1721–1729, <https://doi.org/10.1101/gr.210641.116>.
- [69] F. Meyer, D. Paarmann, M. D'Souza, R. Olson, E.M. Glass, et al., The metagenomics RAST server – a public resource for the automatic phylogenetic and functional analysis of metagenomes, *BMC Bioinformatics* 9 (1) (2008) 386, <https://doi.org/10.1186/1471-2105-9-386>.
- [70] J.G. Caporaso, J. Kuczynski, J. Stombaugh, K. Bittinger, F.D. Bushman, et al., QIIME allows analysis of high-throughput community sequencing data, *Nat. Methods* 7 (5) (2010) 335, <https://doi.org/10.1038/nmeth.1303>.
- [71] E. Plummer, J. Twin, D.M. Bulach, S.M. Garland, S.N. Tabrizi, A comparison of three bioinformatics pipelines for the analysis of preterm gut microbiota using 16S rRNA gene sequencing data, *J. Proteomics Bioinform.* 8 (12) (2015) 283, <https://doi.org/10.4172/jpb.1000381>.
- [72] H. Hasman, D. Saputra, T. Sichert-Ponten, O. Lund, C.A. Svendsen, et al., Rapid whole genome sequencing for the detection and characterization of microorganisms directly from clinical samples, *J. Clin. Microbiol.* 52 (1) (2014) 139–146, <https://doi.org/10.1128/JCM.02452-13>.
- [73] A. Klindworth, E. Pruesse, T. Schweer, J. Peplies, C. Quast, et al., Evaluation of general 16S ribosomal RNA gene PCR primers for classical and next-generation sequencing-based diversity studies, *Nucleic Acids Res.* 41 (1) (2013), <https://doi.org/10.1093/nar/gks808> e1–e1.
- [74] Available online at <http://biopet-docs.readthedocs.io/en/latest/pipelines/flexiprep/>.
- [75] A. Bankevich, S. Nurk, D. Antipov, A.A. Gurevich, M. Dvorkin, et al., Spades: a new genome assembly algorithm and its applications to single-cell sequencing, *J. Comput. Biol.* 19 (5) (2012) 455–477, <https://doi.org/10.1089/cmb.2012.0021>.
- [76] S.Y. Anvar, L. Khachatryan, M. Vermaat, M. Van Galen, I. Pulyakhina, et al., Determining the quality and complexity of next-generation sequencing data without a reference genome, *Genome Biol.* 15 (12) (2014) 555, <https://doi.org/10.1186/s13059-014-0555-3>.
- [77] F. Pedregosa, G. Varoquaux, A. Gramfort, V. Michel, B. Thirion, et al., Scikit-learn: machine learning in python, *J. Mach. Learn. Res.* 12 (October) (2011) 2825–2830.
- [78] M. Burrows, D.J. Wheeler, A Block-Sorting Lossless Data Compression Algorithm, (1994).
- [79] P. Ferragina, G. Manzini, Indexing compressed text, *J. ACM (JACM)*. 52 (4) (2005) 552–581, <https://doi.org/10.1145/1082036.1082039>.
- [80] Centrifuge prebuild indexes: <ftp://ftp.ccb.jhu.edu/pub/infphilo/centrifuge/data>.
- [81] Available online at: <http://www.bioinformatics.babraham.ac.uk/projects/fastqc>.
- [82] M. Martin, Cutadapt removes adapter sequences from high-throughput sequencing reads, *EMBnet J.* 17 (1) (2011) 10, <https://doi.org/10.14806/ej.17.1.200>.
- [83] T. Magoc, S.L. Salzberg, Flash: fast length adjustment of short reads to improve genome assemblies, *Bioinformatics* 27 (21) (2011) 2957–2963, <https://doi.org/10.1093/bioinformatics/btr507>.
- [84] N.A. O'Leary, M.W. Wright, J.R. Brister, S. Ciuffo, D. Haddad, et al., Reference

- sequence (RefSeq) database at NCBI: current status, taxonomic expansion, and functional annotation, *Nucleic Acids Res.* 44 (D1) (2015) D733–D745, <https://doi.org/10.1093/nar/gkv1189>.
- [85] F.P. Breitwieser, S.L. Salzberg, Pavian: interactive analysis of metagenomics data for microbiomics and pathogen identification, *BioRxiv.* (2016) 084715, , <https://doi.org/10.1101/084715>.
- [86] M. Rho, H. Tang, Y. Ye, FragGeneScan: predicting genes in short and error-prone reads, *Nucleic Acids Res.* 38 (20) (2010), <https://doi.org/10.1093/nar/gkq747> e191–e191.
- [87] R.C. Edgar, Search and clustering orders of magnitude faster than BLAST, *Bioinformatics* 26 (19) (2010) 2460–2461, <https://doi.org/10.1093/bioinformatics/btq461>.
- [88] W.J. Kent, BLAT - the blast-like alignment tool, *Genome Res.* 12 (4) (2002) 656–664, <https://doi.org/10.1101/gr.229202>.
- [89] A. Wilke, T. Harrison, J. Wilkening, D. Field, E.M. Glass, et al., The M5nr: a novel non-redundant database containing protein sequences and annotations from multiple sources and associated tools, *BMC Bioinformatics* 13 (1) (2012) 141, <https://doi.org/10.1186/1471-2105-13-141>.
- [90] H.B. Mann, D.R. Whitney, On a test of whether one of two random variables is stochastically larger than the other, *Ann. Math. Stat.* 18 (1) (1947) 50–60, <https://doi.org/10.1214/aoms/1177730491>.
- [91] Y. Benjamini, Y. Hochberg, Controlling the false discovery rate: a practical and powerful approach to multiple testing, *J. R. Stat. Soc. Ser. B* 57 (1) (1995) 289–300.
- [92] C.J. Van Rijsbergen, *Information Retrieval*, second ed., Butterworth-Heinemann, 1979 1979.
- [93] N. Nagarajan, C. Cook, M. Di Bonaventura, et al., Finishing genomes with limited resources: lessons from an ensemble of microbial genomes, *BMC Genomics* 11 (2010) 242, <https://doi.org/10.1186/1471-2164-11-242>.
- [94] D.J. Edwards, K.E. Holt, Beginner's guide to comparative bacterial genome analysis using next-generation sequence data, *Microb. Inform. Exp.* 3 (1) (2013) 2, <https://doi.org/10.1186/2042-5783-3-2>.
- [95] E.A. Grice, J.A. Segre, The skin microbiome, *Nat. Rev. Microbiol.* 9 (4) (2011) 244–253, <https://doi.org/10.1038/nrmicro2537>.