

Crowdsourced-benchmarking of computational pipelines for metagenomic taxonomy profiling – the sbv IMPROVER Microbiomics Challenge

Carine Poussin¹, Nicolas Sierro¹, Vincenzo Belcastro¹, James Battey¹, Stéphanie Boué¹, Fernando Meyer², Elena Scotti¹, Alexandros Petrelis¹, Manuel C Peitsch, Nikolai V Ivanov¹, Julia Hoeng¹

¹Philip Morris International R&D, Philip Morris Products S.A., Quai Jeanrenaud 5, 2000 Neuchâtel, Switzerland (part of Philip Morris International group of companies)

²CAMI, Department of Computational Biology of Infection Research, Helmholtz Centre for Infection Research (HZI), Braunschweig, Germany

The Microbiota Composition Prediction Challenge

BACKGROUND

A growing body of evidence suggests that the equilibrium of microbiome communities and their interaction with their host plays an important role in maintaining host health. **Changes in microbiome communities (dysbiosis)** may be a cause or a consequence of the development of many diseases.

The choice of adequate computational metagenomics methods is **decisive in determining microbial taxonomy profiles and functions accurately**. However, there is no consensus yet about the best analytical and computational approaches to use.

Crowdsourcing initiatives such as Critical Assessment of Metagenome Interpretation (CAMI) (<http://www.cami-challenge.org/>) have started to benchmark algorithms focusing on specific aspects of metagenomics data analysis (1).

The sbv IMPROVER crowdsourcing project (sbvimprover.com), as a means to verify methods and data in systems biology, has already shown its usefulness in benchmarking computational methods through crowdsourcing.

THE CHALLENGE

To build and expand upon what has been done by CAMI, the design of the sbv IMPROVER Microbiomics Challenge focuses on the influence of sample complexity (number of species) and sequence bias (AT vs. GC-rich) on the quantification of microbial communities at various taxonomic levels based on shotgun sequencing data.

GOAL

The challenge aims to assess the performance of **metagenomics computational analysis pipeline(s) as a whole to recover relative abundance and taxonomy assignment of bacterial communities**.

Scientific questions

- Which pipelines best recover bacterial community composition and relative sequence read abundance at phylum, genus and species taxa rank?
- Do technical biases and specific microbial composition affect the performance?

Timelines

The challenge was open from mid-November 2017 to mid-June 2018 (7 months)

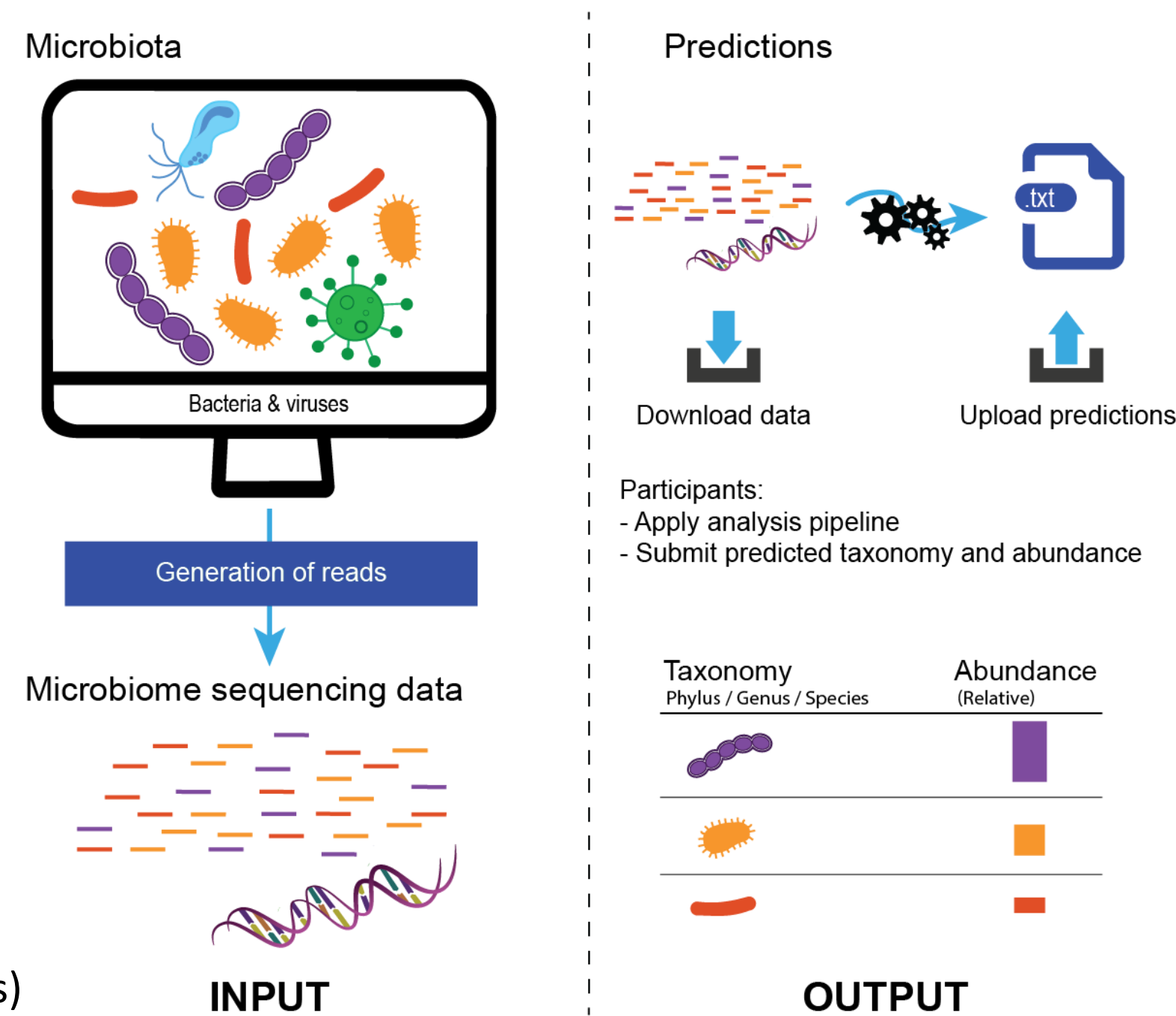
Rules

Eligibility for scoring depended on:

- Submission completeness
- Compliance to data format and challenge rules

Best performers' award

Travel bursary of \$2'000 for a conference (3 best teams)



DATASET

Nineteen paired-end reads samples were provided to the participants (two fastq files per sample). Among these 19 samples, 15 samples have been generated computationally using the ART (2) set of simulation tools using parameter settings to simulate next-generation sequencing reads from an Illumina HiSeq4000 sequencer ("Simulated"), and the four remaining samples correspond to actual Illumina HiSeq4000 sequencing data from a commercial standard DNA control with known bacteria mix ("Real").

Sequencing Reads	# Species	# Samples	Characteristic
Simulated	10	3	unbiased
	40	3	GC-rich bacteria
	120	3	AT-rich bacteria
	500	3	
	1000	3	+ Mouse reads
Real	10*	2 (lib. 1)	DNA control (ZymoBIOMICS)
	10*	2 (lib. 2)	– 2 library preparations

(*8 bacteria species, 2 yeast species)

=19 samples in total

SCORING

Anonymized participants' submissions were split by samples and taxonomic ranks, and evaluated against the corresponding gold standards, i.e., known composition of the respective samples using metrics implemented by CAMI in the OPAL (<https://github.com/CAMI-challenge/OPAL>) software as defined below. Various reference computational pipelines were evaluated by us and CAMI in addition to participants' submissions. These latest predictions, not eligible for the challenge awards, were nevertheless scored for comparison with other approaches tested by challenge participants.

Qualitative/Binary classification metrics: allow to assess **how well a particular method detects the presence or absence of an organism** in comparison to the gold standard.

- F1 score** [0-1]: is the harmonic mean of precision and recall. As it does not put weight on true negatives, which are predominant in microbiome datasets, it is suitable in this context.

Quantitative/Abundance metrics: allow to assess **how well a particular method reconstructs the relative abundances** in comparison to the Gold Standard.

- L1-norm:** distance shows the distance between relative abundance vectors.
- Weighted UniFrac:** distance metric for comparing the composition of microbial communities, which incorporates phylogenetic information (3).

Score aggregation for final ranking: consisted of weighted sum of ranks of each sample-taxa-metric rank per team. The best performer was the team with the lowest aggregated rank.

RESULTS

