

RESOURCE ARTICLE

The impact of contaminants on the accuracy of genome skimming and the effectiveness of exclusion read filters

Eleonora Rachtman¹  | Metin Balaban¹ | Vineet Bafna²  | Siavash Mirarab³ 

¹Bioinformatics and Systems Biology Graduate Program, UC San Diego, CA, USA

²Department of Computer Science and Engineering, UC San Diego, CA, USA

³Department of Electrical and Computer Engineering, UC San Diego, CA, USA

Correspondence

Siavash Mirarab, Department of Electrical and Computer Engineering, UC San Diego, CA 92093, USA.

Email: smirarab@ucsd.edu

Funding information

This work was supported by the National Science Foundation (NSF) grant IIS-1815485 to ER, MB, VB, and SM. Computations were performed on the San Diego Supercomputer Center (SDSC) through the Extreme Science and Engineering Discovery Environment (XSEDE), which is supported by National Science Foundation grant number ACI-1548562.

Abstract

The ability to detect the identity of a sample obtained from its environment is a cornerstone of molecular ecological research. Thanks to the falling price of shotgun sequencing, genome skimming, the acquisition of short reads spread across the genome at low coverage, is emerging as an alternative to traditional barcoding. By obtaining far more data across the whole genome, skimming has the promise to increase the precision of sample identification beyond traditional barcoding while keeping the costs manageable. While methods for assembly-free sample identification based on genome skims are now available, little is known about how these methods react to the presence of DNA from organisms other than the target species. In this paper, we show that the accuracy of distances computed between a pair of genome skims based on k-mer similarity can degrade dramatically if the skims include contaminant reads; i.e., any reads originating from other organisms. We establish a theoretical model of the impact of contamination. We then suggest and evaluate a solution to the contamination problem: Query reads in a genome skim against an extensive database of possible contaminants (e.g., all microbial organisms) and filter out any read that matches. We evaluate the effectiveness of this strategy when implemented using Kraken-II, in detailed analyses. Our results show substantial improvements in accuracy as a result of filtering but also point to limitations, including a need for relatively close matches in the contaminant database.

KEYWORDS

contamination, filtering, genome skimming, Kraken, shotgun sequencing

1 | INTRODUCTION

Anthropogenic pressure and other natural causes have resulted in severe disruption of the global ecosystems in recent years, including loss of biodiversity and invasion of non-native flora and fauna. Conservationists, struggling with an unprecedented rate of extinction, are using innovative approaches to measure the changing biodiversity of the planet. Genome sequencing provides an attractive alternative to physical sampling and cataloging, as falling costs have made it possible to shotgun sequence a reference specimen sample for at most \$10 per Gb (with another \$60 for sample preparation).

However, the analysis typically requires assembling and finishing a reference genome, which can still be prohibitively costly. It could be many decades before the biodiversity of our planet is represented in the form of finished genomes (and cataloged genomic variants) and before biodiversity measurements for each population can be acquired on an ongoing basis.

The standard molecular technique for measuring biodiversity at the organismal level is barcoding (Hebert, Cywinska, Ball, & de-Waard, 2003; Savolainen, Cowan, Vogler, Roderick, & Lane, 2005; Taberlet, Coissac, Pompanon, Brochmann, & Willerslev, 2012), which involves DNA sequencing of taxonomically informative and

group-specific marker genes e.g., mtDNA COI (Hebert et al., 2003; Seifert et al., 2007), 12S/16S (Vences, Thomas, Meijden, Chiari, & Vieites, 2005), plastid genes (Hollingsworth et al., 2009), and ITS (Schoch et al., 2012). Existing reference databases and computational methods enable measurements of biodiversity using barcodes (Ratnasingham & Hebert, 2007; Steinke, Vences, Salzburger, & Meyer, 2005; Taberlet et al., 2012). However, since barcodes are short regions, their phylogenetic signal is limited (Hickerson, Meyer, Moritz, & Hedin, 2006). For example, 896 of the 4,174 species of wasps could not be distinguished from other species using COI barcodes (Quicke et al., 2012).

As an alternative, a genome skim is a low-coverage acquisition of short reads from a sample, typically around 1–5 Gbp (Coissac, Hollingsworth, Lavergne, & Taberlet, 2016; Dodsworth, 2015), providing 0.1–10× coverage, and usually insufficient for assembling nuclear contigs. Falling sequencing costs have made genome-skimming cost-effective while providing richer data than barcoding, but the data is harder to analyze. Skimming applications often rely on assembling organelle genomes (Malé et al., 2014; Weitemier et al., 2014) from their over-represented reads. This approach throws away the vast majority of the reads, potentially limiting the resolution. Moreover, organelle genomes may not represent the rest of the genome and are not always easy to assemble. Ideally, we should use both reads from both nuclear and organelle genomes. However, methods that seek to mine all information from genome skims must be assembly-free and map-free and face additional challenges.

Recently, Sarmashghi, Bohmann, Gilbert, Bafna, and Mirarab (2019) developed a method, Skmer, that accurately computes genomic distance between genome skims by simply analyzing k -mers (short substrings of length k) in both genome skims. Skmer is based on three principles. First, as observed by Ondov et al. (2016), the Jaccard index, J (the size of the intersection of two sets divided by the size of their union) between k -mer sets of the two genomes can be computed efficiently. Second, J can be used to estimate the genomic distance (D) between two species by carefully accounting for dependence on coverage, sequencing error, and genome length. Third, both coverage and error rate can be computed from genome skim data by modelling histograms of k -mer frequencies. By combining these three principles, Skmer provides excellent accuracy in estimating distances between genome skims. These distances can then be used for taxonomic identification and phylogenetic placement (Balaban, Sarmashghi, & Mirarab, 2019) of query genome skims with respect to a set of reference genome skims. Previous results have shown high accuracy and increased resolution compared to barcodes when using genome skims for taxonomic identification (Balaban et al., 2019; Sarmashghi et al., 2019).

The Skmer methodology, however, completely ignores the very real possibility that a genome skim includes extraneous reads originating from other species, often bacteria, virus, or fungi, that cohabit inside the biological organism. With a slight abuse of terminology, we refer to all reads originating from species other than the target species being identified as contamination. Contamination of genome skims is unavoidable in many cases as microorganisms that coexist

with a species are often hard or impossible to separate from the original sample. To make matters worse, laboratory protocols used for genome skimming also can add human and other forms of contamination. The standard organelle-based analyses of skims manage to deal with sequencing errors and contamination by focusing on and assembling a small portion of the reads. These contaminants have the potential to mislead the Jaccard-based calculation of distance using methods such as Skmer. Thus, to take advantage of all reads across the genome, contaminants will have to be dealt with.

In this article, we study the impact of contamination on Skmer estimates of the genomic distance. We then study whether the negative impact of contamination can be reduced using “exclusion filters”: search every read of a skim against a library of all known contaminants (e.g., bacterial, fungal, and viral genomes), filter out reads that map to the library, and use the remaining reads to compute the distance. The efficacy of this exclusion filtering approach is unclear and can depend on several factors, which we thoroughly explore here. We study these effects both based on a theoretical model and in careful simulation and empirical analyses using a leading read matching tool called Kraken-II (Wood, Lu, & Langmead, 2019).

2 | MATERIAL AND METHODS

2.1 | Theoretical exposition

Consider two genomes of equal length and separated by genomic distance D , defined as the portion of positions that do not match in a perfect alignment of the two genomes. Let ρ denote the proportion of k -mers in one species that are absent in the other. The Jaccard-index of k -mers is given by (Figure 1a):

$$J = \frac{\text{Intersection of } k\text{-mer sets}}{\text{Union of } k\text{-mer sets}} = \frac{1-\rho}{2-(1-\rho)} = \frac{1-\rho}{1+\rho}.$$

We model the evolutionary process producing the two genomes as follows. Starting from one genome, mutate each position independently with probability D to get the second genome. With this model, ρ becomes a random variable. Ignoring the dependency between adjacent k -mers and assuming k -mer independence, we get $\mathbb{E}(\rho) = 1 - (1-D)^k$. As Fan, Ives, Surget-Groba, and Cannon (2015) show, we can estimate:

$$\hat{D} = 1 - \left(\frac{2J}{1+J} \right)^{\frac{1}{k}} \quad (1)$$

Skmer further models coverage and sequencing error and uses.

$$\hat{D} = 1 - \left(\frac{2(\zeta_1 L_1 + \zeta_2 L_2) J}{\eta_1 \eta_2 (L_1 + L_2) (1+J)} \right)^{1/k} \quad (2)$$

where η_i , ζ_i , and L_i are parameters related to coverage, error, and genome length, all automatically estimated by Skmer from k -mer profiles.

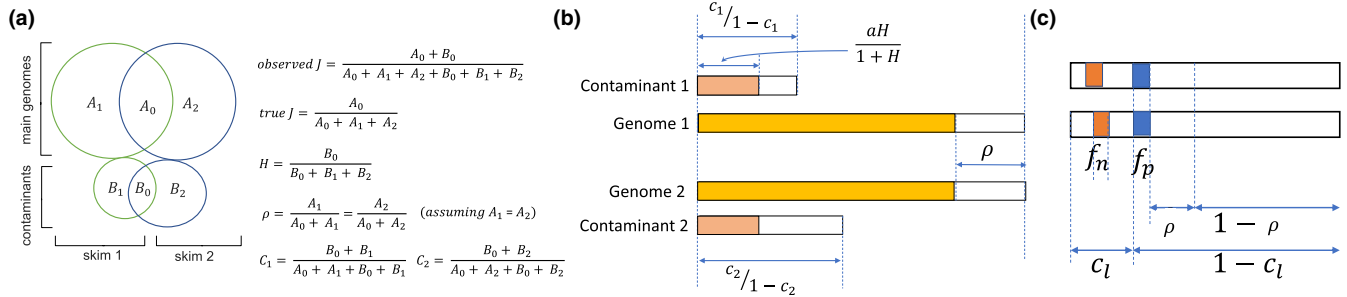


FIGURE 1 Model. (a) Definition of terms. Contaminating k -mers change the estimated Jaccard in a complex manner. (b) Assuming equal lengths for the two genomes, we measure all quantities as a fraction of the number of k -mers in each genome: $1-\rho$ of the k -mers are shared between the base genomes; additionally, out of a total of $a = \frac{c_1}{1-c_1} + \frac{c_2}{1-c_2}$ contaminants, $\frac{aH}{1+H}$ are shared, making the total number of shared k -mers equal to $(1-\rho) + \frac{aH}{1+H}$ and the total number of distinct k -mers in the union equal $(1+\rho) + \frac{a}{1+H}$. See Appendix A1 for details. (c) Impact of false positives and negatives in contaminant removal in the disjoint contaminant scenario. We keep $(1-c_l)(1-f_p) + c_l f_n$ of the k -mers in each set, with the intersection proportion being $(1-c_l)(1-f_p)(1-\rho)$

As the simpler equation is easier to manipulate, we use (1) in our theoretical exposition. However, our empirical results will use the Skmer software, which uses (2). Throughout the paper, we use the relative error to quantify any error in estimating D :

$$\text{relative error of } \hat{D} = \frac{\hat{D} - D}{D} \quad (3)$$

where D is the true genomic distance and $\hat{D}$ is the estimated genomic distance.

2.1.1 | Impact of contamination

Contamination can clearly alter Jaccard and hence the estimated genomic distance (Figure 1a). The impact of contamination depends on factors such as the amount and exact composition of contaminants. For exposition purposes, let us assume that an identical proportion of k -mers (denoted by c_l) of both skims are contaminated, and contaminant k -mers are entirely disjoint between the two genome skims. Then, J becomes a function of c_l :

$$J = \frac{(1-c_l)(1-\rho)}{2-(1-c_l)(1-\rho)} \approx \frac{(1-c_l)(1-D)^k}{2-(1-c_l)(1-D)^k}$$

where the approximation is achieved by replacing ρ with its expectation.

Under these assumptions, Jaccard reduces under contamination, and the extent of reduction depends on c_l and to a lesser degree on D (Figure S1a). If the impact of contamination on Jaccard is ignored, the distance will be overestimated at a level that strongly depends on the true distance (Figure S1b). When D is sufficiently high, substantial levels of contamination result in relatively low errors. However, with smaller distances, contamination can drastically increase the relative error. At $D = 0.1\%$ (e.g., within species differentiation), 3% contamination is enough to cause 100% relative error. Thus, under the simple disjoint contamination model, contamination has a large negative impact *only* when the distance between base genomes is small.

Disjoint contamination assumption, however, is quite strong. When both samples are contaminated with the same species (say, human), the assumption of disjoint contaminant k -mers can mislead. To generalize, consider two genomes with an equal number of k -mers L . Let c_l denote the fraction of the k -mers from sample 1 that are contaminated. Then, the ratio of contaminated k -mers to true k -mers in genome 1 is given by $\frac{c_l}{1-c_l}$ (Appendix A1). Define c_2 in an analogous

fashion for genome 2, and let $a = \frac{c_1}{1-c_1} + \frac{c_2}{1-c_2}$ (Figure 1a). Removing the

disjoint contamination assumption, define the Jaccard index between the k -mers of the contaminants of the two samples as H . Then, as shown in Appendix A1 and Figure 1b,

$$J = \frac{(1-\rho)(1+H) + aH}{(1+\rho)(1+H) + a} \quad (4)$$

Plotting this formula shows that depending on H , the estimated Jaccard may overestimate or underestimate the true Jaccard, and converting the Jaccard to distance without any consideration of contaminants can lead to over or underestimate the true distance (Figure 2a). Once again, error depends on the true distance D , where most dramatic error happens when the distance is low and H is also low. Introduction of H shows that contamination can result in both over and underestimation of error. In particular, for larger values of D , if contaminants are similar between the two samples, relatively low levels of contamination can lead to severe underestimations of distance. For example, with $D = 0.18\%$, if the samples are contaminated at 5% with somewhat similar species with $H = 0.5$, the estimated distance will be underestimated by 43%.

2.1.2 | Impact of exclusion filtering

One approach to deal with contamination is using exclusion filters: search *all* reads in a genome skim against a (potentially incomplete) library of known contaminants and filter out reads that match the

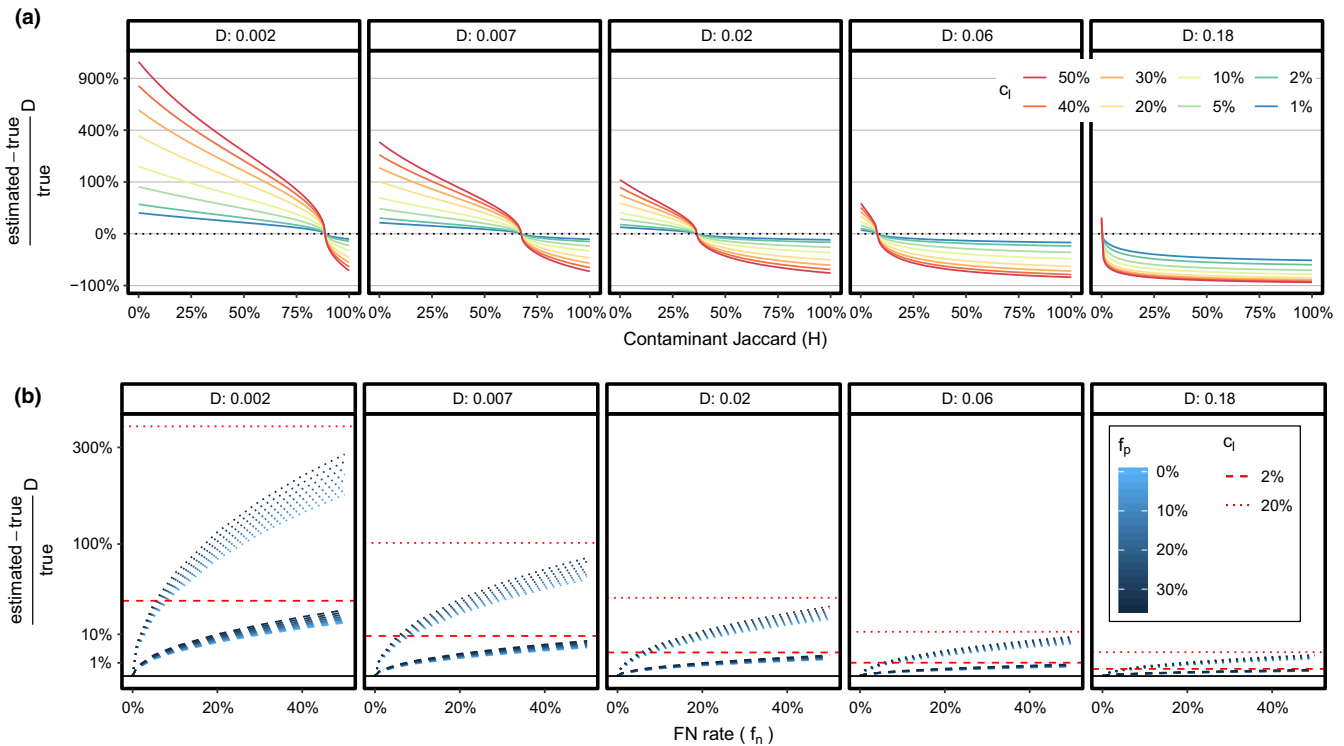


FIGURE 2 Theoretical modeling. (a) Impact of contamination on the genomic distance estimated from Jaccard according to theoretical expectation assuming contaminant k -mers of the two skims have a Jaccard of H (4). For several D and varying H , the relative error is shown for eight contamination levels $c_i = c_1 = c_2$. (b) Error in Skmer distance (computed using (1), with Jaccard approximated using (5)) in the presence of filtering and with the disjoint contaminant k -mer assumption for various levels of FP portion (f_p), FN (f_n) rate, and c_i . Red lines show the error in the absence of filtering. y-axis is in square root scale and $k = 31$

library. This approach will impose a trade-off between two types of possible errors. A false positive (FP) occurs when we incorrectly filter out a read that belongs to the target genome. A false negative (FN) occurs when we fail to filter out a read that belongs to contaminants, perhaps due to an insufficient similarity between the read and genomes included in the exclusion library. The exact choice of the method and parameters used for mapping reads to reference contaminant libraries, in addition to the composition of the reference library, create a trade-off between FP and FN error. The trade-off poses an important question: which type of error, FP or FN, is more damaging? Falling back on the disjoint contaminant k -mer assumption, we can approximate the impact of FP and FN on J given one more assumption: A k -mer shared between the two genome skims is either kept or removed from both skims.

Let f_p be the portion of all k -mers that we remove by mistake (FP) and f_n be the portion of the contaminant k -mers that we fail to remove (FN). The proportion of k -mers shared between genome skims after filtering is $(1 - c_i)(1 - f_p)(1 - \rho)$ (Figure 1c). Additionally, $(1 - c_i)(1 - f_p) + c_i f_n$ of the k -mers in each set are retained after filtering for the total number of unique k -mers to be $2((1 - c_i)(1 - f_p) + c_i f_n) - (1 - c_i)(1 - \rho)(1 - f_p)$. Thus,

$$J = \frac{(1 - c_i)(1 - \rho)(1 - f_p)}{(1 - c_i)(1 + \rho)(1 - f_p) + 2f_n c_i} \quad (5)$$

By plotting this equation as we vary the four parameters (D , c_i , f_p , and f_n), we observe that filtering can successfully reduce the impact of contamination under many but not all conditions (Figure 2b and Figure S2). Filtering can be very effective in making Jaccard index close to what we would obtain without contamination, and overall, Jaccard is more sensitive to FN errors than it is to FP errors. The impact of filtering on genomic distance depends on the level of contamination, false negatives, and most of all, the true genomic distance. Reassuringly, in this model, the estimated accuracy of distance after filtering is reasonably high in most cases (Figures S2 and S3). Nevertheless, in the most challenging cases, filtering cannot sufficiently reduce the error. With $D = 0.2\%$, unless f_n is low or c_i is moderate, error can be very high. Overall, f_p errors are less damaging than f_n . Practically, it seems that with $f_n \leq 0.2$, highly accurate estimates of distance are possible unless contamination levels are very high and the genomic distance is very low.

2.2 | Empirical analyses using Kraken-II

Our empirical experiments validate the effectiveness of exclusion filtering, focusing on a leading k -mer-based read mapping tool called Kraken-II, originally designed for metagenomics and adopted here for contamination filtering. We start by describing Kraken-II and then detail the setup for the four experiments performed.

2.2.1 | Kraken

Kraken-II works by mapping all k -mers of a read to k -mers in a reference library and calls a read a match if the number of k -mers matching is strictly larger than a user-provided threshold called the confidence level (α). Kraken-II uses LCA-mapping to find the lowest taxonomic level at which the read can be confidently matched. It also uses wild-carding s random positions of each k -mer (Brinda, Sykulski, & Kucherov, 2015) to increase the sensitivity of matches. We will explore both k and α settings but fix other parameters. We set minimizer length $l = k$ or use the maximum allowed l value (31) for reference databases built with $k > 31$. We set s , the number of wild-card positions, to its maximum allowable value, $l/4$. We design our reference Kraken-II libraries to include a set of potential contaminants; as query, we use the bag of all reads in a genome skim (details described below). We will use microbial genomes to simulate contamination, and thus, all reference libraries we use are microbial. In contrast, our base genomes are Eukaryotic (plants or insects).

2.2.2 | Experiments

We present four experiments that explore the impact of D (equivalently, ρ), c , f_p , and f_n . In addition, we test the running time of Kraken-II. Below, we describe the setup used in each experiment.

Exploring FN and FP of Kraken-II

We start by examining the sensitivity of Kraken-II to two parameters: k and α . We also consider the completeness of the reference library, which is expected to have a direct effect on f_p and f_n rates. A lack of sufficiently close genomes to the contaminant can prevent Kraken-II from finding a match, and the presence of genomes similar to non-contaminant genomes can cause FP matches. Thus, we define a third variable, M , as the genomic distance between a query and its closest match in the library. We control M by carefully selecting species included in the reference library and those used as query.

To control M , we use an available reference phylogeny of 10,575 bacterial and archaeal genomes (Zhu et al., 2019). Five genomes from this set had IDs that did not exist in NCBI anymore. We assigned remaining genomes to the reference library (10,460 genomes), the query set (100), or both (10). Based on the available phylogeny, we select 10 sets of 10 query genomes such that all genomes in a set have similar patristic (tree) distance to their closest leaf in the tree, not counting the query genomes. These sets had a mean tree distance of (0.01,0.02,0.04,0.06,0.09,0.10,0.18,0.23,0.57,1.20) and at most 25% divergence from the mean. We also randomly chose 10 genomes to be added to both reference and query sets. Then, for each of the 110 query genomes, we used Mash to compute M : its minimum distance to any of the 10,470 reference genomes. We then binned the 110 queries into 10 bins based on M (Table S1). Finally,

we added 10 plant genomes (Table S2) to the set of query genomes in every bin. Plant species are from a different domain of life compared to the reference set and should not match the library; thus, they allow us to measure FP and TN rates.

We built Kraken-II reference libraries for selected k values (ranging from 23 to 35) using the 10,470 bacterial and archaeal reference genomes. Kraken-II only allows adding additional custom genomes of interest to its existing standard reference libraries. We used Kraken-II RefSeq viral genome database as a base library. All custom reference libraries were constructed without masking low complexity sequences.

We used the ART simulation tool (Huang, Li, Myers, & Marth, 2012) with HiSeq 2,500 single read profile, 150 bp read length with 10bp standard deviation to generate ≈ 1.4 GB of synthetic reads (1,000x coverage for each genome) for all query genomes. Every query genome was then downsampled to 1G for normalization purposes.

Reads in each query bin were queried against every constructed reference library for each k using several confidence levels (0–0.3). We then calculated TP, FP, FN, TN for every bin. TP is the count of bacterial/archaeal reads matched to Bacterial or Archaeal domains; FP is the count of plant reads matched to Bacterial or Archaeal domains; TN is the number of plant reads that are left unclassified by Kraken; FN is the number of unclassified bacterial/archaeal sequences. We use standard definitions of FPR = (FP)/(FP + TN) and Recall = TP/(TP + FN) and construct ROC curves in the standard fashion for every tested condition.

Skmer distances (simulation)

We next study the impact of contamination on distances computed from pairs of genome skims simulated from *Drosophila* assemblies. We first emulate the disjoint contaminant scenario by contaminating one of the two genome skims at a level c . We used *D. simulans* w501 to simulate the contaminated genome skim and used *D. simulans* WXD1, *D. sechellia*, or *D. yakuba* to simulate the uncontaminated skim. Based on assemblies, the distances between *D. simulans* w501 and the three other species are 0.2%, 2.1%, and 6.3%, respectively, and we treat these as true distances. To add contamination, we use the same 110 query genomes described earlier but bin M into four ranges: [0,0], (0,0.05], (0.05,0.15], and (0.15,0.25], which include 10, 43, 19, and 17 species, respectively, corresponding to a total size of 37 Mb, 76 Mb, 37 Mb, and 35 Mb. Since our base *Drosophila* genomes are roughly 150 Mb in size, we can add up to 25% contaminant reads for all bins, except for the (0,0.05] bin, where we can add up to 60%. We concatenated all the genomes in each bin and used ART with the same settings indicated above to generate contaminant reads, which we then mixed with reads simulated from the main genome at levels varying from 0% to 60% (for the second bin) or to 25% (for all other bins) for a total of 0.1Gb per skim (thus, no more than 1x coverage). These read contamination levels translate to similar k -mer contamination levels (Table S3). We report the relative error in estimated distances as we increase the contamination

level, both with and without Kraken-II filtering. Kraken is run with the same reference library used in the previous analysis.

We then simulate a scenario where both genome skims are contaminated with overlapping sets of species. Here, we only use the $M \in (0, 0.05]$ bin and fix read contamination level to 15%. To control H , we randomly split bacterial reads into three parts: two unique parts and one part that served as an overlap. Every sample was generated by mixing unique and overlap contaminant portions with *Drosophila* genome skims at controlled ratios, with overlap set to 0%–50%. Since unique parts can have evolutionary similar species, even the case of 0% overlap results in some k -mer overlap. Thus, we estimated contamination overlap (H) empirically using Jellyfish (Marçais & Kingsford, 2011) and saw it varied between 11% and 41% (Table S4). Finally, to have $H=0\%$, we added the disjoint set experiment with $M \in (0, 0.05]$ and $c_i=15\%$ to this set as well.

Skmer distances on real data

To move beyond simulations, we also evaluate effectiveness on real data with real contaminants. To do so, we utilized data from a recent *Drosophila* assembly study by Miller, Staber, Zeitlinger, and Hawley (2018). We subsampled available short-read sequencing data (e.g., SRA files) to obtain 100 Mb genome skims for 14 *Drosophila* species. We removed adapters, deduplicated and merged paired-end reads using BBtools (Bushnell, Rood, & Singer, 2017). Then, we determined distances for all pairs of genomes before and after filtering them with Kraken-II. Distance error for every pair of genomes was estimated relative to the true distance defined to be the value computed by running Skmer on corresponding assemblies. In this experiment, we used a standard reference library available from Kraken-II distribution. This database includes RefSeq assemblies of all available bacterial, archaeal, viral and human (GRCh38) genomes as well as the UniVec_Core subset of the UniVec database (a total of 168,483 genomes, as of July 2019). We used default Kraken-II settings.

Impact of filtering on the phylogeny

On the real *Drosophila* data, we also infer phylogenetic trees from distances and measure phylogenetic error. To estimate the phylogeny, we use FASTME 2.0 software (Lefort, Desper, & Gascuel, 2015) with JC69+ Γ (Jin & Nei, 1990) model of evolution. Alpha parameter of JC69+ Γ model is set equal to 1, which is the default value in FastME. We infer phylogenies from distance matrices obtained from assemblies and from genome skims before and after filtering. As the gold standard reference tree used for error calculations, we use the tree obtained from Open Tree of Life (OTL) (Hinchliff et al., 2015, Figure S4). We estimated branch lengths of the true tree using OTL tree topology and assembly distances under JC69+ Γ model. We measure phylogenetic error using three metrics. (1) Normalized Robinson and Foulds (1981) (RF) distance is the total number of branches not matching. (2) Normalized weighted RF (wRF) distance is similar, but each present or absent branch in each tree is weighted

by the absolute difference between its lengths in the two trees, and then the total sum is normalized by the sum of branch lengths of the two trees. (3) Fitch-Margoliash (Fitch & Margoliash, 1967) is the weighted least squares error (FME) for species i , given as: $Q(i) = \sum_{j \neq i} (D_{ij}/d_{ij} - 1)^2$ where D_{ij} is the (corrected) distance between

species i and j , and d_{ij} is sum of the branch lengths on the path connecting i and j on the phylogeny inferred using D . We also report

cumulative FME of a phylogeny, which is $Q = \sum_{i=1}^N Q(i)$. Denoting FME

error on true and estimated phylogenies with $Q(i)$ and $\hat{Q}(i)$, respectively, relative FME error is defined similarly to (3).

3 | RESULTS

3.1 | Sensitivity of Kraken-II (FN and FP analysis)

The ability of the default version of Kraken-II ($k=35, \alpha=0$) to find a match in the database is a direct function of M , the distance of the query to the closest match (Figure S5 and Figure 3a). When the query has a close match in the library (e.g., $D < 0.05$), Kraken-II is able to match 80%–100% of reads, which would result in tolerable f_n rates of 20% or less. As M increases, the ability of Kraken-II to classify degrades linearly with M up until around $M \approx 0.3$ where Kraken-II fails to classify almost all reads (Figure S5). Interestingly, when Kraken-II finds a match, it is often able to classify the read all the way down to the species level (Figure S6).

Consistent with these results, when mixed plant/microbe skims are queried using the default Kraken, the recall of the filtering step is reasonably high (e.g., >85%) and the FP is low (4.5%) for $M \leq 0.05$ (Figure 3a). When $0.05 < M \leq 0.1$ or $0.1 < M \leq 0.15$, there is a substantial reduction in the recall to 67% and 56%, respectively; for $0.15 < M$, recall is less than 33%, and thus filtering is not effective in those conditions.

Given the low recall in some conditions and our expectation that FP error is less damaging than FN, one may aspire to increase the sensitivity of Kraken-II by adjusting its parameters k and α . However, our careful analysis of FP vs. FN shows very limited ability to control the rates in a reasonable range (Figure 3a, b and Figure S7). Many settings of k and α result in FP error above 50% and often close to 100%. Many of the settings also have high FP without improving recall compared to default settings (Figure 3b). The only settings that seem to provide a reasonable trade-off between FPR and recall are $k \in \{35, 32, 28\}$ and $\alpha \leq 0.05$. Focusing on these settings (Figure 3a), we observe that setting $k=28$ and $\alpha=0$ provides a substantial increase in recall but increases FPR to unacceptable levels (55%). $k=32$ improves recall compared to the default setting for $M \geq 0.05$ bins by a consistent but relatively small margin (5%–8%), but also increases the FPR to 8.5%. Overall, changing parameters do not result in substantial improvements over the default settings.

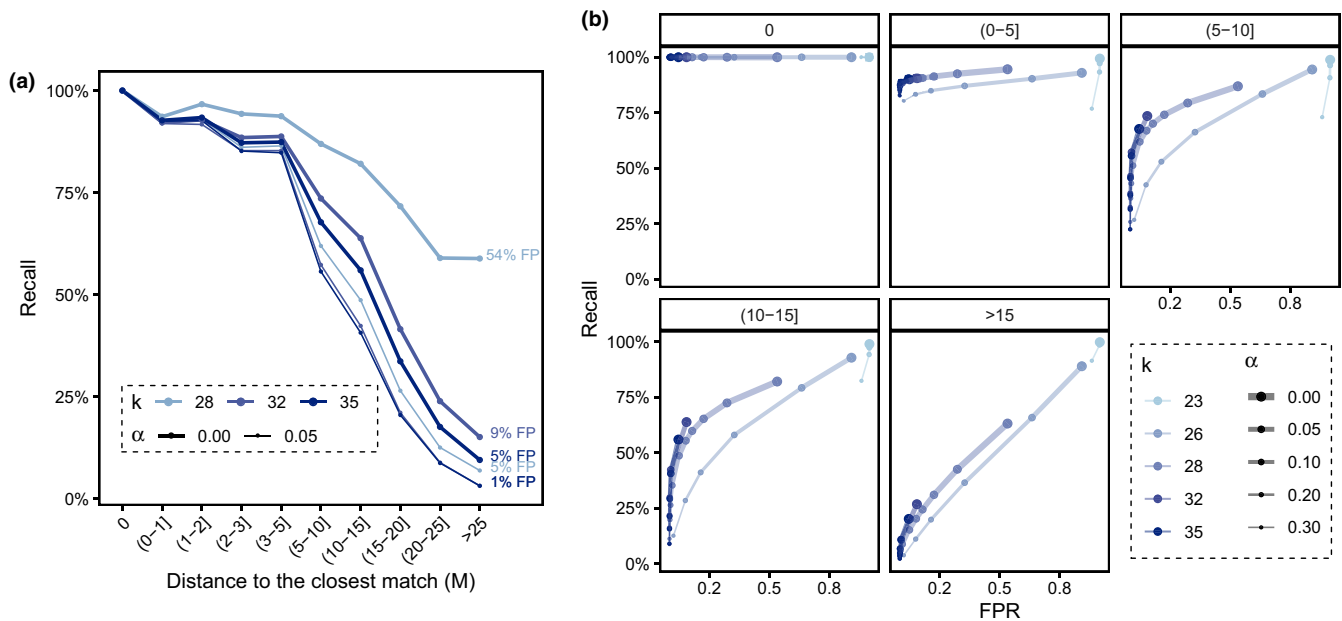


FIGURE 3 Sensitivity analysis of Kraken-II. (a) For three selected values of k and two selected values of α , the lines show recall for different bins of M (x-axis). Each line is labelled with its associated FPR. Note that FPR is a function of the plant genomes, which are identical across bins; thus, FPR is not a function of the match parameter M . (b) ROC curves for all k and α values across different bins of M , reduced to five ranges for ease of visualization (boxes). All k were tested with $\alpha \in \{0.0, 0.05, 0.1, 0.15, 0.2, 0.3\}$. Additionally, $k = 28$ was tested with $\alpha \in \{0.01, 0.02, 0.03, 0.04\}$. See Figure S7 for all 10 bins

3.2 | Impact of filtering on Skmer distances (simulated contaminants)

3.2.1 | Disjoint contaminants

Focusing on simulated contamination between pairs of *Drosophila* genome skims, when only one species is contaminated, increasing the contamination level results in increasing error in estimated Skmer distances, going up to 90% error for $D=2\%$ and 1,000% for $D=0.2\%$ when $c_i=60\%$ (Figure 4a). As theory suggested, here, the strongest detrimental effect appears for $D=0.2\%$.

Filtering using default Kraken-II dramatically reduces the error when the contaminant has an exact or close match in the reference library (Figure 4a). For $M \leq 0.05$, remarkably high levels of contamination are tolerated after filtering. For example, for $0 < M \leq 0.05$ and $D=2.1\%$, even with high c_i in 25%–50%, distances have only 0.3%–4% relative error after filtering. For $D=6.3\%$, error after filtering is never more than 5% for $M \leq 0.05$. Even in the most challenging case of $D=0.2\%$, $c_i=25\%$ leads to only a 6% error after filtering in contrast to a 206% error before filtering. Despite the improved accuracy overall, in some cases, filtering can increase the error slightly but noticeably, perhaps due to the FP filtering of correct reads. For $M=0$ and $D=6.3\%$, if contamination is below 5%, no filtering is better than filtering, which always results in $\approx 0.6\%$ relative error regardless of the level of contamination. Interestingly, in some cases, filtering can result in underestimation of distances (e.g., up to 1% for $D=2.1\%$ and $M=0$).

In contrast, for contaminants without a close match with $M > 0.05$, filtering fails to fully remove contaminants. Nevertheless, for $0.5 < M \leq 0.15$, filtering has substantial benefits. For example,

the error is reduced from 180% and 16% with no filtering to 65% and 6%, respectively for $D=0.2\%$ and $D=2.1\%$. These reductions, while substantial, may not be sufficient. Even worse, for $M > 0.15$, filtering has very little or no ability to reduce the error and decreases or increases the error by very small margins.

Finally, changing k, α settings of Kraken-II does not consistently improve the accuracy above and beyond the default setting (Figure S8). Using $\alpha=0.05$ can very slightly reduce the error for the $D=6.3\%$ case but is not dramatically different. Thus, we will exclusively use the defaults in the next experiments.

3.2.2 | Overlapping contaminants

When both skims are contaminated with overlapping species, as theory suggested, we see underestimation of distances (Figure 4b). These underestimations can be dramatic, going all the way down to -100% (i.e., the estimated distance is 0). Once again, filtering using Kraken is able to improve results dramatically, resulting in a relative error that does not exceed 23% for $D=0.2\%$ and is at most 6% in the remaining cases.

3.3 | Impact of filtering on Skmer distances (real contaminants)

In the experiment on real unassembled *Drosophila* sequences, absent any filtering, Skmer often underestimates distances (Figure 5a). The underestimation of distances is consistent with

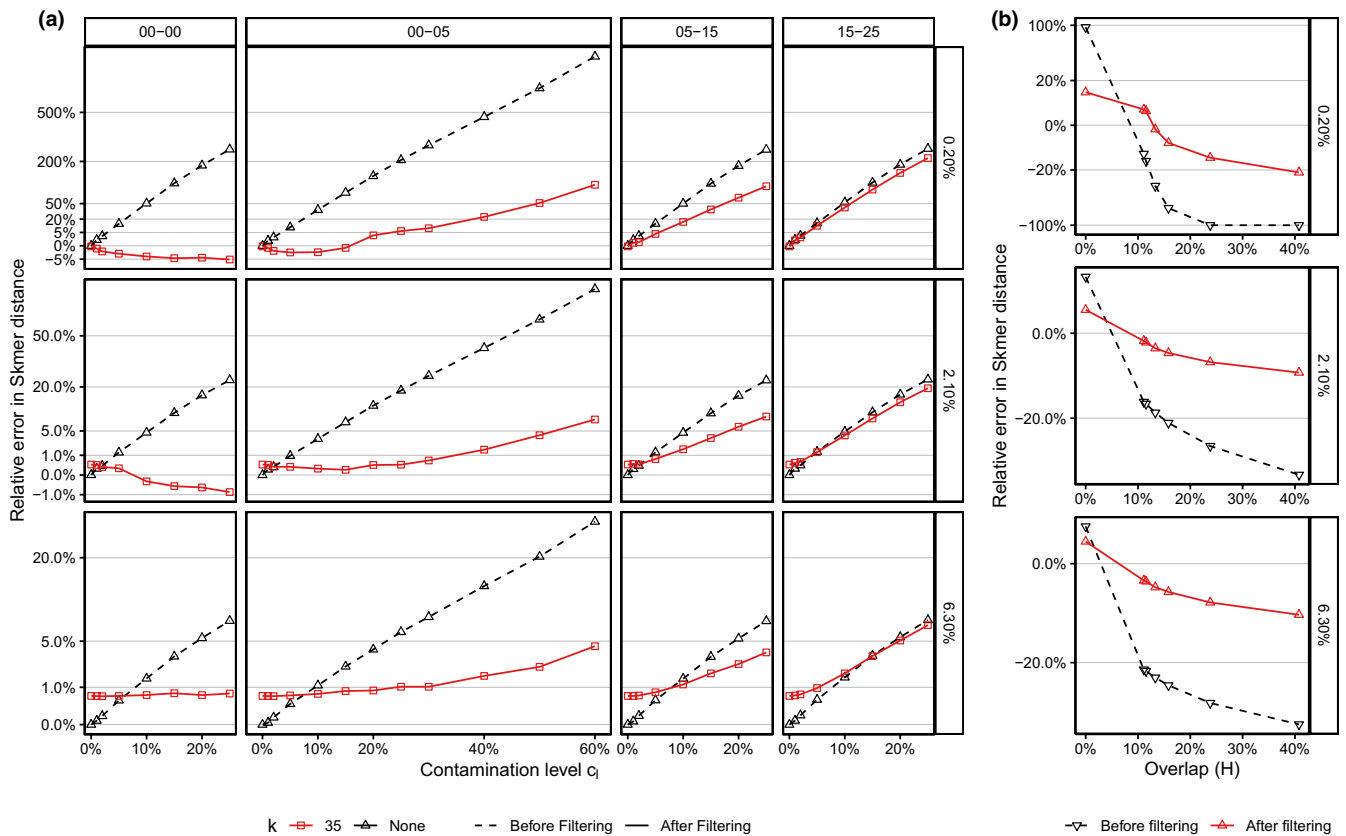


FIGURE 4 Filtering on simulated *Drosophila* genome skims. Relative error of Skmer distances without (dashed) and with (solid) Kraken-II filtering. Three pairs of *Drosophila* are chosen to be at true distance $D = 0.2\%$, $D = 2.1\%$, or $D = 6.3\%$ (rows). Contaminants are selected such that they are at distance M from their closest match in the reference contaminant library. (a) Simulating contaminants in only one of the two species (disjoint contaminants) for four ranges of M (columns) and various levels of contamination (x-axis). (b) Contaminating both genomes such that the overlap between contaminants measured by Jaccard similarity is $0\% \leq H \leq 41\%$. Here, $c_i = 15\%$ per species and $M \in (0, 0.05]$. Y-axis is on square root scale; see 1.0 for normal scale and a range of k and α values

our theory assuming $H > 0$ (Figure 2). Kraken-II run on these data identifies between 5.5% and 15.1% of the reads as belonging to human or microbes (Figure 5b). Interestingly, for most *Drosophila* species, Kraken-II assigns ~40%–50% of the matched reads to one of three genera (*Homo*, *Acetobacter*, and *Clostridium*), indicating that many pairs of genome skims have similar contaminants (i.e., $H > 0$). Therefore, the underestimation of distances matches the theory. Consistent with this explanation, we observe that the error in computed distances is associated with the percentage of the reads found by Kraken-II to be of human or microbial origin (Figure S9).

Filtering reads using Kraken-II dramatically reduces the errors in Skmer distances (Figure 5c). Over all pairs, the mean absolute relative error reduces from 9.1% before filtering to 3.4% after filtering. In some cases, reductions are dramatic. For example, the relative error in pairwise distances between *D. virilis* and *D. bipectinata*, *D. eugracilis* and *D. mauritiana*, decreased from 46.2%, 36.9% and 35.9% before filtering to 1.3%, 0.8%, and 1.0% after filtering. In a minority of cases, error increased after filtering, but the increase in error never exceeded 8% (*D. mauritiana* vs. *D. mojavensis*) while reductions in error could be as high as 45% (*D. virilis* vs. *D. bipectinata*)

(Figure 5c). The wide range of error reductions is unsurprising given that the actual level of contamination in the original sample can vary substantially. The magnitude of improvement in distance estimates is positively correlated with the percentage of reads filtered (Figure 5c), the genomic distance (Figure S10a), and the magnitude of the error before filtering (Figure S10b).

When there is error after Kraken-II filtering, it tends to be due to overestimation of distance, as opposed to underestimation observed before filtering, suggesting that Kraken-II perhaps over-filters reads. Most extreme cases of over-filtering involve distance estimates between a single species, *D. mojavensis*, and other species such as *D. bipectinata*, *D. mauritiana* and *D. virilis*. The *D. mojavensis* is the only species with high levels of Kraken-II filtering but low error rates in pairwise comparisons. Interestingly, *D. mojavensis* also includes the highest levels of contamination from unknown sources.

3.4 | Impact on phylogenetic reconstruction

The phylogeny inferred from Skmer distances computed from the assembly and modelled using JC69+ Γ is topologically identical to

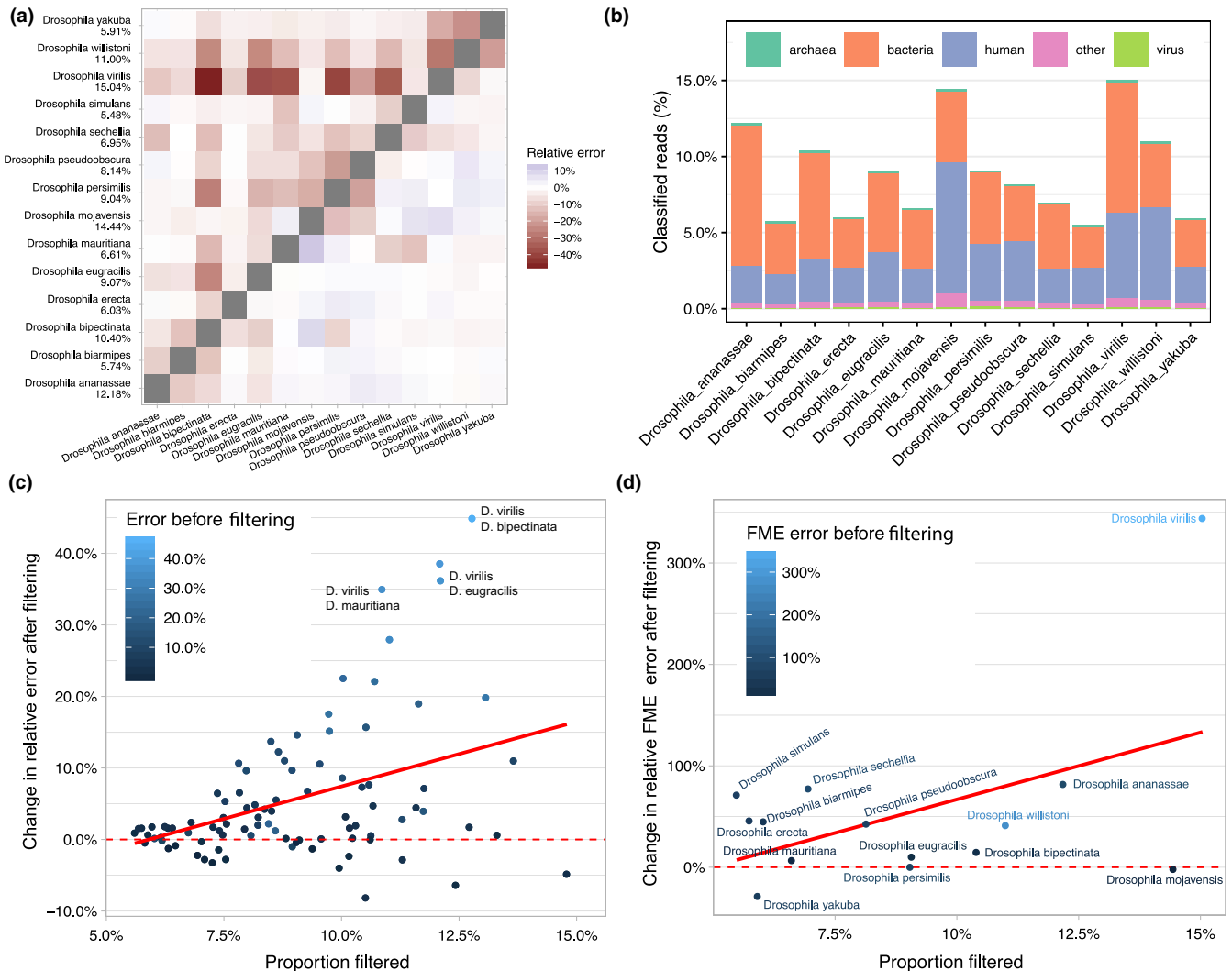


FIGURE 5 Filtering of real contaminants. (a) Relative distance error before (upper triangle) and after (lower triangle) filtering per pair of *Drosophila* species. Numeric labels on the y-axis represent the percentage of reads filtered per species. (b) Percent of reads classified by Kraken-II to different groups. “Other” corresponds “cellular organisms” (shared between domains). (c) Change in the relative distance error after filtering. Positive values indicate a reduction in error. Highlighted dots correspond to *Drosophila* pairs mentioned in the text. (d) Change in relative FME error per species after filtering. Solid red line: a trend line fitted to the points

the gold standard OTL phylogeny, and its total FME error is only 0.03 (Figure S4). However, the phylogeny estimated using the same method but using genome skims has two wrong branches (RF = 4), and a FME of 1.26. Thus, absent filtering, genome skims produce trees with substantial error.

Improvements in estimated genomic distances due to filtering translate to improved phylogenetic trees. The tree topology improves only slightly and has one incorrect branch (RF = 2) after filtering. However, the improvements in estimated branch lengths, as reflected in wRF and total FME error, are dramatic. Filtering leads to nearly 70% decrease in total FME metric, from 1.26 to 0.38, and a similar level of reduction is observed for wRF (Table S5). Examining individual branch lengths, the phylogeny using filtered data is much more similar to the true tree (Figure S4).

When we use FME to measure the impact of filtering on the phylogenetic error of individual species, we observe patterns consistent

with reductions in distance error (Figure 5d). Individually, majority of species have reduced FME after filtering, with the most extreme FME reduction happening for *D. virilis* by nearly 350%. Consistent with previous results, we observe that the FME error of *D. mojavensis* does not decrease (but it also does not increase). As expected, gains in phylogenetic error measured by FME correlate with the amount of filtering performed by Kraken (Figure 5d). However, the correlation disappears when the species *D. virilis* is excluded from the analysis (Figure S11).

3.5 | Running time

We assessed running time performance of Kraken-II on skims of five randomly selected species from different domains of life (Table S6) using both single- and multithreaded (24 threads) modes of

operation. Run time was found to linearly increase with the skim size, regardless of the number of threads used (Figure S12). With 1 Gb of reads, the running time of the single-threaded version was below 100 s on an Intel Xeon CPU in all cases we tested. Running Kraken-II with 24 threads reduced the speed by a factor of 10, offering a significant improvement. The main limitation with Kraken-II is its significant memory requirement during queries, which requires between 100 Gb and 120 Gb for our reference libraries.

4 | DISCUSSION

The use of genome skimming in the literature has mostly relied on assembled organelle genomes (e.g., Coissac et al., 2016; Dodsworth, 2015; Malé et al., 2014; Weitemier et al., 2014). These approaches rely on assembly construction pipelines (e.g., Jin et al., 2019) to remove contaminants (i.e., to avoid misassembly or to filter out misassembled contigs). Elsewhere, we have advocated going beyond organelle genomes and using all reads in an assembly-free fashion to increase the resolution of taxonomic identification (Balaban et al., 2019; Sarmashghi et al., 2019). However, this goal has been hampered by the presence of contaminants. This study showed a relatively effective way of dealing with contamination, hence bringing genome skimming based on nuclear reads one step closer to reality.

Our study showed that Kraken-II is able to find contaminants that are within 5%–10% genomic distance to the closest match in a reference library in a computationally efficient manner. Our modelling showed that FP errors were perhaps less detrimental to distance calculations than FN. Analysis of different k and α parameters did not reveal parameter combinations that could improve upon the using default settings of Kraken-II ($k=35, \alpha=0$). Analyses of real data demonstrated that contamination removal can dramatically improve Skmer distance estimates in the presence of contaminants. These more accurate distance metrics computed after filtering can lead to reduced phylogenetic branch length error by up to 70% and can also improve the tree topology.

4.1 | Usefulness of theoretical models

Simplified assumptions allowed us to establish theoretical models of the impact of contamination on estimated distance. The theory predicted that even small levels of similarity between contaminants (H) can lead to substantial underestimation of distance when the distance is large. Consistently, on the real data, where distances are often >0.10 , we observe underestimation by 5% or more in 47 out of 91 pairs. Our results also showed high levels of similarity between contaminants of *Drosophila* genomes (where three genera made up 40%–50% of contaminants). Thus, there is a reassuring match between the theoretical model and the observed data.

Kraken filtering improved accuracy on simulated and real data. On real data, it occasionally over-corrected errors, leading to overestimation of the distance. These may be due to FP filtering, reduced

coverage after filtering, or other factors not fully understood here. In our runs of Kraken-II on real data, we observed 5%–15% filtering. The lower value can be explained by ~5% Kraken-II FP rate when run under its default setting. The upper value is consistent with ~10% contamination level, a scenario that can happen in real sequencing projects (Merchant, Wood, & Salzberg, 2014; Sangiovanni, Granata, Thind, & Guarracino, 2019).

Another potential use of the theory could have been developing filter-free methods of dealing with contamination. Just as the impact of coverage and error on Jaccard can be modelled, we can compute the Jaccard index with no filtering but *correct* for the modelled impact of contamination on Jaccard. Given reliable estimates of H , c_1 , and c_2 , we can manipulate (4) to update (1) and obtain:

$$\hat{D} = 1 - \left(\frac{(2+a)J}{1+J} - \frac{aH}{1+H} \right)^{1/k}$$

Adding coverage and error models, we can update (2) to:

$$\hat{D} = 1 - \left(\frac{(2+a)(\zeta_1 L_1 + \zeta_2 L_2)J}{\eta_1 \eta_2 (L_1 + L_2)(1+J)} - \frac{aH}{1+H} \right)^{1/k} \quad (6)$$

This equation allows for filter-free contamination-aware distance calculation. Unfortunately, however, this equation is extremely sensitive to correct estimation of all parameters, including H , c_1 , and c_2 (Figure S13). Even small mistakes (1%–5% relative error) in the estimated contamination level or Jaccard can lead to dramatic errors in the estimated distance computed using (6). Since the computation of these parameters is noisy, we do not advocate this filter-free method despite its theoretical elegance.

4.2 | Filtering methods

Filtering requires a tool to answer queries of the following type: “Does this particular read belong to one of the genomes in a given reference library?”. We chose Kraken-II for answering these queries because of its high accuracy and reasonable scalability, as established in several bench-marking studies from the metagenomics field (McIntyre et al., 2017; Meyer et al., 2019; Sczyrba et al., 2017; Ye, Siddle, Park, & Sabeti, 2019). It is also one of the most widely-used tools, with active user support and stable and robust software development. Further, we explored three parameters of Kraken-II: k -mer length, confidence score, and database content. Other parameters such as minimizer length (mostly relevant to storage and not accuracy) and minimizer space count are not explored here (Brinda et al., 2015). In our experiments, we kept the number of wild-carding positions at its recommended upper limit and turned masking off, but there might be a set of settings which, in combination with masking, can produce a more optimal sensitivity. We note that results from Wood et al. (2019) have indicated that Kraken-II is not very sensitive to particular parameter settings.

Alternatives to Kraken-II exist, and future studies can compare them to Kraken-II for genome skimming. BLAST (Altschul, Gish,

Miller, Myers, & Lipman, 1990) and MegaBLAST (Morgulis et al., 2008) are the obvious alternatives but are overkill for our problem. These tools perform alignment and can yield higher sensitivity than Kraken-II but are orders of magnitude slower (Wood & Salzberg, 2014; Ye et al., 2019). However, they produce more precise results (maps to individual species) than what we need.

Beyond alignment tools, most alternatives to Kraken-II are also *k*-mer-based but differ in the way the reference library is constructed and how the query is run. *k*-mer-based methods include LMAT (Ames et al., 2013), and CLARK(-S) (Ounit & Lonardi, 2016; Ounit, Wanamaker, Close, & Lonardi, 2015). Benchmarking studies (e.g., McIntyre et al., 2017; Meyer et al., 2019; Sczyrba et al., 2017; Ye et al., 2019) do not indicate any consistent advantage in using these methods over Kraken-II, and many of them are slower. Among sufficiently fast tools are KrakenUniq (Breitwieser, Baker, & Salzberg, 2018), Bracken (Lu, Breitwieser, Thielen, & Salzberg, 2017), and Centrifuge (Kim, Song, Breitwieser, & Salzberg, 2016). KrakenUniq is recommended for use in cases where FP can be detrimental (e.g., in pathogen identification/diagnoses), but our theory and empirical data suggest FP is less important and FN in our application. Bracken (Lu et al., 2017), an extension of Kraken-II, is focused on improving aggregated abundance profiles, a feature that is irrelevant to our usage. Centrifuge (Kim et al., 2016) uses FM-index lookups and within-species compression for mapping a read to one or more species. Compared to Kraken-II, Centrifuge is slower and needs more time for building its reference database. We leave its comparison to Kraken for future work.

A separate set of *k*-mer-based methods have been developed for finding RNAseq experiments that include a specific *k*-mer. Solomon and Kingsford (2016) introduced Sequence Bloom Tree (SBT) to allow very fast queries of a *k*-mer versus a reference set of experiments by creating a hierarchy of compressed bloom filters that store *k*-mers. Mantis (Pandey et al., 2018) is an alternative to Bloom filters based on counting quotient filters and is reported to be more memory efficient and faster than SBT-based methods. While these tools have been developed mainly for RNAseq analyses, in the future, they can perhaps be adopted for mapping reads to genomes with minimal changes to the algorithm. In fact, Kraken might implement counting quotient filter data structure in its future releases (Wood et al., 2019).

Beyond these tools, many other metagenomic methods have been designed for finding the taxonomic composition of a mixed sample (e.g., Liu, Gibbons, Ghodsi, & Pop, 2010; Milanese et al., 2019; Nguyen, Mirarab, Liu, Pop, & Warnow, 2014; Segata et al., 2012). However, these tools do not seek to classify every read from anywhere in the genome; they are either marker-based or use composition data. Thus, these tools are irrelevant to our queries.

4.3 | Remaining gaps

In our study, we focused solely on prokaryotic and human contamination. Real contamination is more complex and can include eukaryotic

microorganisms, traces of endosymbionts and diet, and various forms of lab contamination. Thus, many applications will benefit from more inclusive Kraken-II contaminant libraries. At a minimum, fungi need to be considered, especially for plants. Moreover, removing reads from organelle genomes, which are expected to be over-represented, may further improve accuracy.

Luckily, Kraken-II enables a straightforward mechanism for extending reference libraries. Our future efforts will include building a larger library of potential contaminants that includes fungi and perhaps expected sources of diet. However, such libraries will have to be group-specific; for example, for skimming insects, we can treat plants as contaminants, whereas in skimming plants, we should treat insects as contaminants. Ideally, individual genome skimming reference libraries for a target group (e.g., all insects) should be furnished with a relevant contaminant library specially designed for that group based on the knowledge of taxonomic groups expected to be present in its diet and its endosymbiont. Clearly, this approach runs into its limitations when endosymbionts or the diet happen to be from species with similar genomes to the target species.

The fundamental limitation of our exclusion filtering approach is that we need to know what broad group of species is expected to contaminate. This limitation is a result of our implicit assumption that a read is correct unless we find evidence to the contrary. Even when such biological knowledge is available – it may not be – this approach can fail to capture lab-introduced contamination (e.g., a plant species that was contaminated with fish due to failures in sample preparation or sequencing on the same lane).

Another issue is the inclusion of the human genome in the reference libraries to find human contamination. When the target species is a mammalian species, reads that belong to the target may incorrectly map to the human (false positive). For example, in a test analysis, we observed that Kraken-II maps roughly 20% of reads from rhino to human. Luckily, this problem has a simple solution. We can require a higher α when mapping to human. Thus, all reads can be searched against the library with the default $\alpha=0$, and those reads that map to human can be mapped again with a high α ; reads are classified as human contamination if they continue to map at the higher α . We tested this approach on the rhino genome and observed that with $\alpha=0.5$, only 2% of reads map to human. We have provided a script for performing this two-step filtering.

Inclusion filtering is an attractive alternative to exclusion filters. Given a reference database of purified (perhaps using exclusion filters) genome skims, we can build a Kraken reference library from species in the skimming reference library. Then, for every new query genome skim, we can use that library to find reads that seem to match the broad taxonomic group of interest and only include those reads in the calculation of Jaccard. Our results indicate that this method would work only if the skimming reference database is so dense that each new query skim is expected to have a close match (e.g., < 5%) to one of the reference skims. Moreover, this approach is predicated on the reference library being free of contaminants. Despite these shortcomings, we believe this approach should be further explored in the future.

Finally, better algorithms for read matching seem necessary. Our results showed that Kraken-II provides a reasonable solution. Nevertheless, the method remains incapable of finding domain level matches when the closest match is moderately distant from the query. We believe it is possible to design more sensitive read mapping techniques that can match a species even when its closest match is > 10% distance. Note that in genome skimming, we are only interested to know whether a read belongs to a large taxonomic group, as opposed to metagenomics, when abundances and exact matches are desired. Given the less demanding needs of the skimming application, we anticipate that better algorithms can be developed in the future to increase recall with little or no loss of specificity and speed.

ACKNOWLEDGEMENTS

We thank Aryn P. Wilder for valuable suggestions on filtering from mammalian genomes.

AUTHOR CONTRIBUTIONS

All authors conceived the idea. S.M. and V.B. developed a theoretical model. E.R. implemented the pipeline and performed experiments. M.B. completed phylogenetic experiments. All authors contributed to the analyses of data and the writing. All authors read and approved the final manuscript.

DATA AVAILABILITY STATEMENT

Scripts and summary data tables are publicly available on https://github.com/norarachtkraken_scripts.git. Raw data used in the manuscript is deposited in https://github.com/norarachtkraken_raw_data.git. Additionally, data repositories are stored in zenodo <https://doi.org/10.5281/zenodo.3588625> (Rachtman, Balaban, Bafna, & Mirarab, 2019a) and <https://doi.org/10.5281/zenodo.3588569> (Rachtman, Balaban, Bafna, & Mirarab, 2019b). The detailed description of genomic datasets used in our experiments, accession numbers of the assemblies and the exact commands used to simulate genome skims are provided in Appendix S1.

ORCID

Eleonora Rachtman  <https://orcid.org/0000-0002-6104-5750>

Vineet Bafna  <https://orcid.org/0000-0002-5810-6241>

Siavash Mirarab  <https://orcid.org/0000-0001-5410-1518>

REFERENCES

- Altschul, S. F., Gish, W., Miller, W., Myers, E. W., & Lipman, D. J. (1990). Basic local alignment search tool. *Journal of Molecular Biology*, 215(3), 403–410.
- Ames, S. K., Hysom, D. A., Gardner, S. N., Lloyd, G. S., Gokhale, M. B., & Allen, J. E. (2013). Scalable metagenomic taxonomy classification using a reference genome database. *Bioinformatics*, 29(18), 2253–2260.
- Balaban, M., Sarmashghi, S., & Mirarab, S. (2019). APPLES: Scalable distance-based phylogenetic placement with or without alignments. *Systematic Biology*, <https://doi.org/10.1093/sysbio/syz063>.
- Breitwieser, F. P., Baker, D. N., & Salzberg, S. L. (2018). KRAKENUNIQ: Confident and fast metagenomics classification using unique k-mer counts. *Genome Biology*, 19(1), 198.
- Brinda, K., Sykulska, M., & Kucherov, G. (2015). Spaced seeds improve k-mer-based metagenomic classification. *Bioinformatics (Oxford, England)*, 31(22), 3584–3592.
- Bushnell, B., Rood, J., & Singer, E. (2017). BBMerge – Accurate paired shotgun read merging via overlap. *PLoS ONE*, 12(10), 1–15.
- Coissac, E., Hollingsworth, P. M., Laverne, S., & Taberlet, P. (2016). From barcodes to genomes: Extending the concept of DNA barcoding. *Molecular Ecology*, 25(7), 1423–1428.
- Dodsworth, S. (2015). Genome skimming for next-generation biodiversity analysis. *Trends in Plant Science*, 20(9), 525–527.
- Fan, H., Ives, A. R., Surget-Groba, Y., & Cannon, C. H. (2015). An assembly and alignment-free method of phylogeny reconstruction from next-generation sequencing data. *BMC Genomics*, 16(1), 522.
- Fitch, W. M., & Margoliash, E. (1967). Construction of phylogenetic trees. *Science*, 155(3760), 279–284.
- Hebert, P. D. N., Cywinska, A., Ball, S. L., & deWaard, J. R. (2003). Biological identifications through DNA barcodes. *Proceedings of the Royal Society B: Biological Sciences*, 270(1512), 313–321.
- Hickerson, M. J., Meyer, C. P., Moritz, C., & Hedin, M. (2006). DNA Barcoding Will Often Fail to Discover New Animal Species over Broad Parameter Space. *Systematic Biology*, 55(5), 729–739.
- Hinchliff, C. E., Smith, S. A., Allman, J. F., Burleigh, J. G., Chaudhary, R., Coghill, L. M., ... Cranston, K. A. (2015). Synthesis of phylogeny and taxonomy into a comprehensive tree of life. *Proceedings of the National Academy of Sciences*, 112(41), 12764–12769.
- Hollingsworth, P. M., Forrest, L. L., Spouge, J. L., Hajibabaei, M., Ratnasingham, S., van der Bank, M., ... Little, D. P. (2009). A DNA barcode for land plants. *Proceedings of the National Academy of Sciences*, 106(31), 12794–12797.
- Huang, W., Li, L., Myers, J. R., & Marth, G. T. (2012). ART: A next-generation sequencing read simulator. *Bioinformatics*, 28(4), 593–594.
- Jin, J.-J., Yu, W.-B., Yang, J.-B., Song, Y., DePamphilis, C. W., Yi, T.-S., & Li, D.-Z. (2019). GetOrganelle: a fast and versatile toolkit for accurate de novo assembly of organelle genomes. *bioRxiv*. <https://doi.org/10.1101/256479>
- Jin, L., & Nei, M. (1990). Limitations of the evolutionary parsimony method of phylogenetic analysis. *Molecular Biology and Evolution*, 7(1), 82–102.
- Kim, D., Song, L., Breitwieser, F. P., & Salzberg, S. L. (2016). Centrifuge: Rapid and sensitive classification of metagenomic sequences. *Genome Research*, 26(12), 1721–1729.
- Lefort, V., Desper, R., & Gascuel, O. (2015). FASTME 2.0: A comprehensive, accurate, and fast distance-based phylogeny inference program. *Molecular Biology and Evolution*, 32(10), 2798–2800.
- Liu, B., Gibbons, T., Ghodsi, M., & Pop, M. (2010). METAPHYL: Taxonomic profiling for metagenomic sequences. In *Proceedings - 2010 IEEE International Conference on Bioinformatics and Biomedicine, BIBM 2010* (pp. 95–100). Hong Kong, China: IEEE. <https://ieeexplore.ieee.org/document/5706544>
- Lu, J., Breitwieser, F. P., Thielen, P., & Salzberg, S. L. (2017). Bracken: Estimating species abundance in metagenomics data. *PeerJ Computer Science*, 3, e104.
- Malé, P.-J.-G., Bardon, L., Besnard, G., Coissac, E., Delsuc, F., Engel, J., ... Chave, J. (2014). Genome skimming by shotgun sequencing helps resolve the phylogeny of a pantropical tree family. *Molecular Ecology Resources*, 14(5), 966–975.
- Marçais, G., & Kingsford, C. (2011). A fast, lock-free approach for efficient parallel counting of occurrences of k-mers. *Bioinformatics*, 27(6), 764–770.
- McIntyre, A. B. R., Ounit, R., Afshinnekoo, E., Prill, R. J., Hénaff, E., Alexander, N., ... Mason, C. E. (2017). Comprehensive benchmarking and ensemble approaches for metagenomic classifiers. *Genome Biology*, 18(1), 182.

- Merchant, S., Wood, D. E., & Salzberg, S. L. (2014). Unexpected cross-species contamination in genome sequencing projects. *PeerJ*, 2, e675.
- Meyer, F., Bremges, A., Belmann, P., Janssen, S., McHardy, A. C., & Koslicki, D. (2019). Assessing taxonomic metagenome profilers with OPAL. *Genome Biology*, 20(1), 51.
- Milanese, A., Mende, D. R., Paoli, L., Salazar, G., Ruscheweyh, H.-J., Cuenca, M., ... Sunagawa, S. (2019). Microbial abundance, activity and population genomic profiling with mOTUs2. *Nature Communications*, 10(1), 1014.
- Miller, D. E., Staber, C., Zeitlinger, J., & Hawley, R. S. (2018). Highly contiguous genome assemblies of 15 drosophila species generated using nanopore sequencing. *G3: Genes, Genomes, Genetics*, 8(10), 3131–3141.
- Morgulis, A., Coulouris, G., Raytselis, Y., Madden, T. L., Agarwala, R., & Schaffer, A. A. (2008). Database indexing for production MegaBLAST searches. *Bioinformatics (Oxford, England)*, 24(16), 1757–1764.
- Nguyen, N., Mirarab, S., Liu, B., Pop, M., & Warnow, T. (2014). TIPP: Taxonomic identification and phylogenetic profiling. *Bioinformatics*, 30(24), 3548–3555.
- Ondov, B. D., Treangen, T. J., Melsted, P., Mallonee, A. B., Bergman, N. H., Koren, S., & Phillippy, A. M. (2016). Mash: Fast genome and metagenome distance estimation using MinHash. *Genome Biology*, 17(1), 132.
- Ounit, R., & Lonardi, S. (2016). Higher classification sensitivity of short metagenomic reads with CLARK-S. *Bioinformatics (Oxford, England)*, 32(24), 3823–3825.
- Ounit, R., Wanamaker, S., Close, T. J., & Lonardi, S. (2015). CLARK: Fast and accurate classification of metagenomic and genomic sequences using discriminative k-mers. *BMC Genomics*, 16, 236.
- Pandey, P., Almodaresi, F., Bender, M. A., Ferdman, M., Johnson, R., & Patro, R. (2018). MANTIS: A fast, small, and exact large-scale sequence-search index. *Cell Systems*, 7(2), 201–207.e4.
- Quicke, D. L. J., Alex Smith, M., Janzen, D. H., Hallwachs, W., Fernandez-Triana, J., Laurence, N. M., ... Vogler, A. P. (2012). Utility of the DNA barcoding gene fragment for parasitic wasp phylogeny (Hymenoptera: Ichneumonoidea): Data release and new measure of taxonomic congruence. *Molecular Ecology Resources*, 12(4), 676–685.
- Rachtman, E., Balaban, M., Bafna, V., & Mirarab, S. (2019a). kraken_raw_data_v2. <https://doi.org/10.5281/zenodo.3588625>
- Rachtman, E., Balaban, M., Bafna, V., & Mirarab, S. (2019b). kraken_scripts_v3. <https://doi.org/10.5281/zenodo.3588569>
- Ratnasingham, S., & Hebert, P. D. N. (2007). BOLD: The Barcode of Life Data System (www.barcodinglife.org). *Molecular Ecology Notes*, 7, 355–364.
- Robinson, D., & Foulds, L. (1981). Comparison of phylogenetic trees. *Mathematical Biosciences*, 53(1–2), 131–147.
- Sangiovanni, M., Granata, I., Thind, A. S., & Guarracino, M. R. (2019). From trash to treasure: Detecting unexpected contamination in unmapped NGS data. *BMC Bioinformatics*, 20(4), 168.
- Sarmashghi, S., Bohmann, K. P., Gilbert, M. T., Bafna, V., & Mirarab, S. (2019). Skmer: Assembly-free and alignment-free sample identification using genome skims. *Genome Biology*, 20(1), 34.
- Savolainen, V., Cowan, R. S., Vogler, A. P., Roderick, G. K., & Lane, R. (2005). Towards writing the encyclopaedia of life: An introduction to DNA barcoding. *Philosophical Transactions of the Royal Society B: Biological Sciences*, 360(1462), 1805–1811.
- Schoch, C. L., Seifert, K. A., Huhndorf, S., Robert, V., Spouge, J. L., Levesque, C. A., ... Schindel, D. (2012). Nuclear ribosomal internal transcribed spacer (ITS) region as a universal DNA barcode marker for Fungi. *Proceedings of the National Academy of Sciences*, 109(16), 6241–6246.
- Szczyrba, A., Hofmann, P., Belmann, P., Koslicki, D., Janssen, S., Dröge, J., ... McHardy, A. C. (2017). Critical Assessment of Metagenome Interpretation benchmark of metagenomics software. *Nature Methods*, 14(11), 1063–1071.
- Segata, N., Waldron, L., Ballarini, A., Narasimhan, V., Jousson, O., & Huttenhower, C. (2012). Metagenomic microbial community profiling using unique clade-specific marker genes. *Nature Methods*, 9(8), 811–814.
- Seifert, K. A., Samson, R. A., deWaard, J. R., Houbraken, J., Levesque, C. A., Moncalvo, J.-M., ... Hebert, P. D. N. (2007). Prospects for fungus identification using CO1 DNA barcodes, with *Penicillium* as a test case. *Proceedings of the National Academy of Sciences*, 104(10), 3901–3906.
- Solomon, B., & Kingsford, C. (2016). Fast search of thousands of short-read sequencing experiments. *Nature Biotechnology*, 34(3), 300–302.
- Steinke, D., Vences, M., Salzburger, W., & Meyer, A. (2005). TAXI: A software tool for DNA barcoding using distance methods. *Philosophical Transactions of the Royal Society B: Biological Sciences*, 360(1462), 1975–1980.
- Taberlet, P., Coissac, E., Pompanon, F., Brochmann, C., & Willerslev, E. (2012). Towards next-generation biodiversity assessment using DNA metabarcoding. *Molecular Ecology*, 21(8), 2045–2050.
- Vences, M., Thomas, M., van der Meijden, A., Chiari, Y., & Vieites, D. R. (2005). Comparative performance of the 16S rRNA gene in DNA barcoding of amphibians. *Frontiers in Zoology*, 2, 5.
- Weitemier, K., Straub, S. C. K., Cronn, R. C., Fishbein, M., Schmickl, R., McDonnell, A., & Liston, A. (2014). Hyb-Seq: Combining target enrichment and genome skimming for plant phylogenomics. *Applications in Plant Sciences*, 2(9), 1400042.
- Wood, D. E., Lu, J., & Langmead, B. (2019). Improved metagenomic analysis with Kraken 2. *Genome Biology*, 20(1), 1–16. <https://doi.org/10.1186/s13059-019-1891-0>
- Wood, D. E., & Salzberg, S. L. (2014). Kraken: Ultrafast metagenomic sequence classification using exact alignments. *Genome Biology*, 15(3), R46. <https://doi.org/10.1186/gb-2014-15-3-r46>
- Ye, S. H., Siddle, K. J., Park, D. J., & Sabeti, P. C. (2019). Benchmarking Metagenomics Tools for Taxonomic Classification. *Cell*, 178(4), 779–794.
- Zhu, Q., Mai, U., Pfeiffer, W., Janssen, S., Asnicar, F., Sanders, J. G., ... Knight, R. (2019). Phylogenomics of 10,575 genomes reveals evolutionary proximity between domains Bacteria and Archaea. *Nature Communications*, 10(1), 5477.

SUPPORTING INFORMATION

Additional supporting information may be found online in the Supporting Information section.

How to cite this article: Rachtman E, Balaban M, Bafna V, Mirarab S. The impact of contaminants on the accuracy of genome skimming and the effectiveness of exclusion read filters. *Mol Ecol Resour*. 2020;00:1–13. <https://doi.org/10.1111/1755-0998.13135>