

INVITED REVIEW OPEN ACCESS

What We Talk About When We Talk About Microbial Species

Apurva Narechania¹  | Shyam Gopalakrishnan¹ | M. Thomas P. Gilbert^{1,2} ¹Center for Evolutionary Hologenomics, the Globe Institute, University of Copenhagen, Copenhagen, Denmark | ²University Museum, NTNU, Trondheim, Norway**Correspondence:** Apurva Narechania (narechan@gmail.com) | M. Thomas P. Gilbert (tgilbert@sund.ku.dk)**Received:** 10 August 2025 | **Revised:** 17 October 2025 | **Accepted:** 28 October 2025**Handling Editor:** Joaana Freeland**Funding:** The authors received no specific funding for this work.**Keywords:** compression | information theory | pangenomes

ABSTRACT

Genome annotation, alignment and phylogenetics are at the centre of most studies in evolutionary genomics. These techniques function best when rooted in prior work. Genes are mined from new genomes using evidence from old gene models. These genomes are aligned to well-worn references to create matrices for tree reconstruction. Trees are often populated with well-characterised genomes to add context to the newly sequenced. Genome inference traces a line back to model organisms, yoking the analysis of new genomes to layers of previous knowledge. Here, we present an alternative approach that uses unannotated and unaligned sequence to understand the information diversity of sequence ensembles. Any set of genomes can comprise our sequence ensemble. In a pandemic context, a sequence ensemble might be clinically isolated strains from one day. In a systematic context, a sequence ensemble could be the pangenome available for a clade. The normal bioinformatics playbook would have us align. But we instead compress. A sequence ensemble that compresses easily contains lower information diversity. For pandemics, we can use curves of information diversity to trace genomic novelty and monitor selective sweeps in new strains. For systematics, we can calculate compressibility quickly across all known bacterial taxa, levelling the criteria for species across clades. If we tolerate data loss, we can go one step further and capture structural evolution as we compress. Our approach sacrifices a lot. We skip many of the products of modern bioinformatics like variation anchored to known genes or genome alignment to prescribed references or pangenome graphs. But we gain speed, breadth and the ability to rapidly respond to novelty.

1 | Introduction (The Problem)

Compression encodes information into reduced representations. Whether bits are eliminated through statistical redundancy (lossless compression), or shed entirely (lossy compression), compressed data always have a smaller footprint than the original. The act of compression—its difficulty or ease—communicates information about the original data source. Highly redundant data with many common patterns will compress easily. In contrast, novelty or surprise with little repeated context is difficult

to compress. Evolution creates ensembles of sequences. These ensembles can be represented as pangenomes. Pangenomes are compressible entities, but how compressible depends on pangenome complexity, which we argue can be a proxy for evolutionary strategy.

Genomics is a retrospective field. Existing bioinformatic techniques often model new genomes on sequences annotated in the past (Guigó 2023). Alignment to these reference genomes circumscribes our knowledge of diversity. Large swaths of the

This is an open access article under the terms of the [Creative Commons Attribution-NonCommercial](https://creativecommons.org/licenses/by-nc/4.0/) License, which permits use, distribution and reproduction in any medium, provided the original work is properly cited and is not used for commercial purposes.

© 2025 The Author(s). *Molecular Ecology* published by John Wiley & Sons Ltd.

tree of life are presumably unknown (Wu et al. 2009). For example, much of the sequence from environmental samples passes through annotation filters as undefined (Thomas and Segata 2019). In a read-streaming era (Erich 2015), we need forward-looking techniques that flag genomic novelty by dispensing with references, annotation or alignment. Standard methods are ill-equipped for these volumes. New species are not easily caught in the sparse web of the known.

As genomics has swept through biology, systematics has come to favour molecular character sets to help delimit species boundaries (Baker and DeSalle 1997; Rokas et al. 2003). While morphology is still important, and holdouts have been more than vocal (Neumann et al. 2021), phylogenomics has more recently carried the day. Phylogenomics extends the handful of marker genes that were the foundation of early molecular systematics to matrices that concatenate thousands of orthologues (Zhu et al. 2019). This character explosion has been a boon to systematics, but gene annotation is still anchored to the known.

These orthologous genes are rarely evenly distributed among the genomes that describe a species (Shi and Falkowski 2008). The complete set is one definition of the pangenome, and its complexity was originally defined as the rate of gene accumulation with newly sequenced genomes (Tettelin et al. 2005). Genes found universally comprise a genomic core and are considered indispensable for basic species functions. Genes found sporadically may contribute to strain success in particular niches but may not be essential to their overall biology. The ratio of core genomes to accessory genomes informs genome fluidity (Kislyuk et al. 2011). Species whose genomes are mostly core have closed, less fluid pangenomes. Species with a large fraction of accessory genes are considered open and more fluid. This gene-centric view of orthologues is blind to the diversity in the non-coding genome (The ENCODE Project Consortium 2012). Whole genome alignment to annotated, chromosomal references (Li et al. 2009; Li and Durbin 2009) makes variation in the non-coding genome accessible but again circumscribes its characterisation. If all we know is a linear reference on a single coordinate system, our understanding of the non-coding genome will be limited to what will stick.

More recent pangenome methods attempt to enhance the reference by conveying it as a graph (Eizenga et al. 2020). For example, a species graph through elements of the genomic core would collapse into a single consensus, punctuated by bubbles that code small-scale variation like single nucleotide polymorphisms and small insertion/deletion elements. In contrast, the accessory genome forks the pangenome graph along entirely disparate paths. Graph-based methods attempt to incorporate nuance and novelty into a more complex reference structure. But the game is still the same: new data are aligned to a set of old genomes bound together into a complex, branching network.

Is there another way? Can we measure some other property of whole genomes that isn't contingent on their alignment? Can we de-centre the gene so we aren't limited to the protein-coding genome? Can we dispense with phylogenomics so we aren't spending CPU years deciphering a bifurcating set of species relationships that convey a mere shadow of a more reticulate truth (Doolittle 1999)?

Here, we propose several new information theoretic techniques that reimagine genomes as ensembles of information, containers subject to compression. This view of genomic information does not require annotation. Because we aren't concerned with genes or the contiguous arrangements of genomic elements, we also forgo alignment. We instead describe pangenomes with summary statistics of string-based intersections. In this article, we argue that compression can enhance existing comparative genomic strategies, highlight structural evolution through controlled information loss and democratise the bacterial species question by applying a uniform mathematics across the Linnean taxonomy.

2 | The Toolkit (Entropy)

Our approach is guided by two foundational concepts at the very root of information theory: entropy and relative entropy. Both ideas rely heavily on Claude Shannon's seminal ideas on information introduced in 'A Mathematical Theory of Communication', the founding document of information theory (Shannon 1948). Information is data that reduce uncertainty. Shannon's original formulation resembled the thermodynamic construction of entropy devised for statistical mechanics (Boltzmann 1877). We measure information entropy as

$$H = - \sum_{i=1}^N p_i \ln p_i$$

where N is the set of all possible states, i , and p_i is the probability of the i -th state. This expression quantifies data into bits (base two logarithm) or nats (natural logarithm). The bit is the most irreducible unit of information. A bit is gained when a binary variable is assigned either a 1 or a 0.

In genomics, our data come in sequences. We can measure the entropy of sequences by digesting them into substrings of specific size. In the bioinformatics literature, substrings of biological readouts (DNA, RNA, protein) are called k-mers. In a comparative setting, we are most interested in the entropy of a group of sequences, or a sequence ensemble (Bialek et al. 2001; Crutchfield and Feldman 2003). For genome sequence ensembles, alignment has been the tool of choice. But alignment is computationally arduous and breaks down with evolutionary distance. Fields as diverse as linguistics (Bentz and Alikaniotis 2016), neurobiology (Bialek et al. 1991) and statistical mechanics (Jaynes 1957) have successfully employed entropy to quantify ensemble complexity. In each of these fields, researchers code a linear string of observations and divide it into subsequences, calculating the entropy of each set across the ensemble. In Figure 1, we show how the entropy of genome sequence (e.g., DNA/RNA) typically increases with increasing k-mer size. This is a block entropy curve.

Block entropy curves contain information about the complexity of the ensemble (Crutchfield and Feldman 2003). Systems with more ensemble structure—repeated elements across sequences—will peak at lower entropy. More novelty across sequences yields higher entropy. In genomics, closed pangenomes, with a large core shared across all species genomes, have low entropy. Auxiliary genes unique to subsets

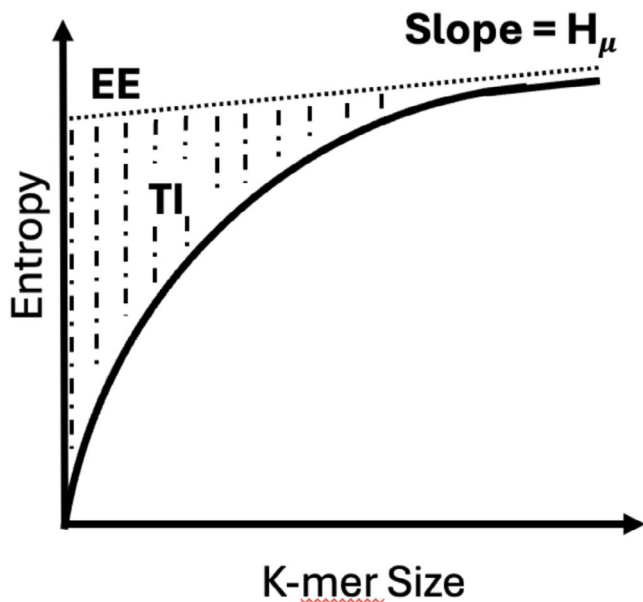


FIGURE 1 | Block entropy curve. We show that entropy increases with k-mer size. We use this curve to calculate source entropy (H_μ), excess entropy (EE) and transient information (TI)⁴¹.

of genomes add entropy to the ensemble system. The uneven distribution of these elements is the hallmark of an open pangenome. But to measure complexity we don't need the annotated and aligned genes. Signal is preserved in unaligned and unannotated k-mers.

Block entropy curves asymptote at the minimum block size required to efficiently capture information diversity across the sequences. We use three quantities calculated from these curves to describe the complexity of a pangenome: source entropy, excess entropy and transient information (Crutchfield and Feldman 2003). The source entropy (H_μ) is the irreducible randomness that remains even as larger block sizes capture most ensemble correlations. H_μ is a direct measure of randomness. Random distributions are hard to compress. A high source entropy is associated with the accumulation of unevenly distributed accessory genes, resulting in a more complex pangenome. The excess entropy (EE) is the non-random fraction of the total information in the system. It is the information we model from redundancies across the ensemble. Alignment is anchored to these redundancies. In fact, alignment works only if enough of these redundancies are spread across the query genomes. Finally, the transient information (TI) measures how much information we must invest to learn H_μ and EE. In Figure 1 we show TI as the area between the block entropy curve and the line defining H_μ . Species with closed pangenomes typically have a lower TI than those open to accumulating gene diversity. Closed pangenomes with a large core set of genes compress at lower k-mer sizes, approaching their H_μ quickly.

3 | More Tools in the Toolkit (The Information Bottleneck)

Entropy is the workhorse of lossless compression. In fact, it defines lossless compression's limit. We cannot compress any further than the entropy of the source. In our context, the block

entropy curve follows compression limits along a k-mer spectrum. Lossless compression preserves all data, but sacrifice can feature evolutionary events by isolating patterns from genomic noise. Using lossy compression, we can identify the core genome of any species without alignment or annotation. Moreover, as we compress we unlock the homologous and non-homologous recombination events that violate vertical signal (Narechania, Bobo, DeSalle, et al. 2025).

To understand how we can detect structural evolution without annotation or alignment, we leverage Shannon's ideas on lossy compression. Shannon based his theory in communication. A sender passes a message to a receiver through a channel. The fundamental problem of communication is reconstructing that message. Communication channels suffer distortion. Data rarely reach the receiver whole. Information entropy represents the limit on how efficiently a message can be compressed in the noise-free ideal.

No channel is noise-less. But the distortion introduced by noisy channels does not doom message passing. A sender can compensate for noise by encoding more information into a message, or a receiver can tolerate some level of distortion while ascertaining a sender's core meaning. Shannon formalised this concept as rate-distortion theory (Shannon 1948). On the sender side, the rate is measured as bits of information per symbol. The sender's message is distorted as it passes through the channel. The sender's rate and the receiver's distortion are inversely related. The function describing the two variables for any given channel informs lossy compression. How much information loss can we tolerate in reconstructing a sender's message? This idea is central not only to information theory and lossy coding, but also to modern machine learning methods that use variational autoencoders to populate the compressive layers of a neural net (Kingma and Welling 2013; Krizhevsky et al. 2017). The two key questions are 1. How well does a dataset compress and 2. How much data can we afford to lose?

We use these concepts to further understand pangenome complexity. Imagine the compression regime in Figure 2. A set of genomes comprising a sequence ensemble is digested into k-mers and compressed into a set number of clusters. This compression is analogous to a communication channel. The more clusters we model, the higher the rate, and the lower the distortion. With fewer clusters, we force the k-mers through a narrower channel and suffer more distortion. If we hold the channel constant and model the same number of clusters across species ensembles, open pangenomes will suffer more data loss than their closed counterparts. Open pangenomes have more complex information to communicate. This is the information bottleneck (Tishby et al. 2000), an idea first proposed in the Natural Language Processing literature. The bottleneck modulates loss in our compression framework. Moreover, the clusters we glean comprise a model of structural evolution. The largest cluster usually represents the core genome. K-mers from recombination regions populate the others. For example, we have shown that we can detect the horizontal gene transfer events that fuelled the evolution of pathogenicity in *Staphylococcus aureus* Clonal Complex 8 (CC8). We do this *de novo*, without knowing anything about the elements themselves. In CC8, the staphylococcal chromosomal cassette

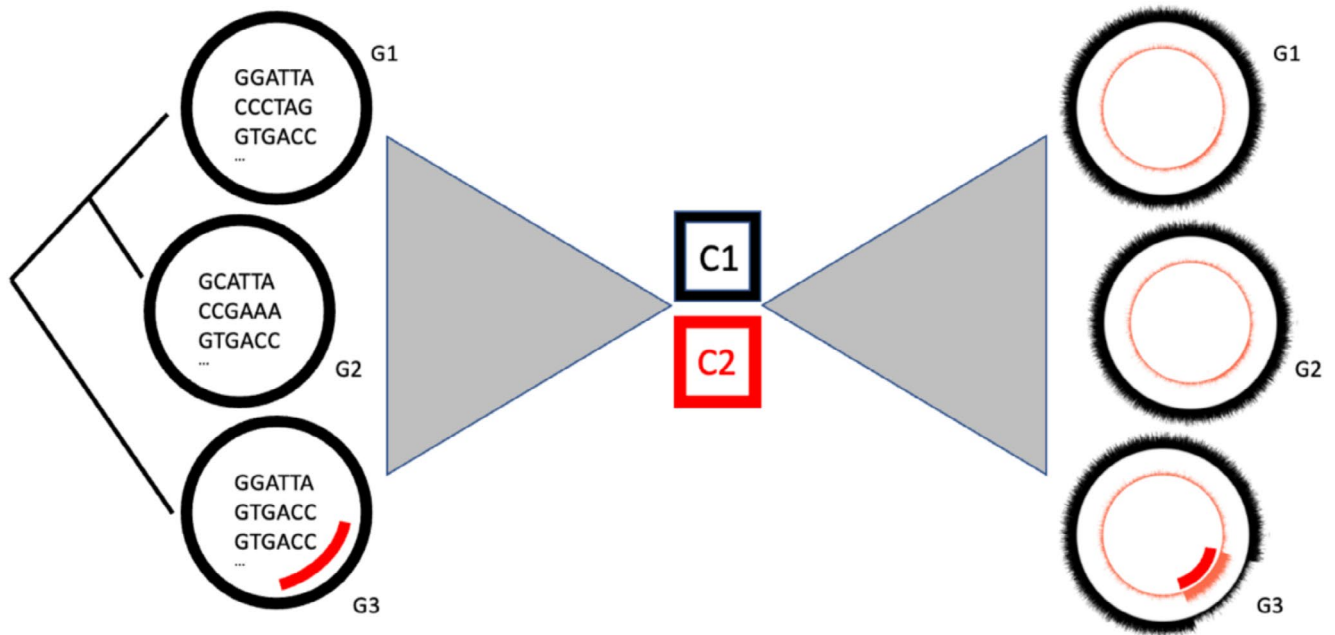


FIGURE 2 | The information bottleneck. As k-mers from our input genomes (G1, G2, and G3) are compressed through a narrow channel, patterns of structural evolution emerge from the resulting clusters (C1 and C2)⁴⁴.

carrying the methicillin resistance gene expanded, accumulating other disease-associated elements as the pathogen moved between North and South America (Narechania, Bobo, DeSalle, et al. 2025). The act of compression therefore learns real biological events without the need to annotate genes, align genomes, build sets of orthologues, or calculate any trees.

4 | Even More Tools in the Toolkit (Relative Entropy)

Block entropy curves measure the compressibility of any sequence ensemble. The bottleneck compresses information into clusters that communicate only the most salient bits. Compressibility is directly related to pangenome composition. Some species are open to genomic input; others have narrower, closed pangenomes. But as we've described it here, entropy treats the entire pangenome distribution as a single entity. This allows us to measure overall complexity but doesn't account for each genome's departure from that distribution. Relative entropy, a measure of how one distribution (any given single genome) diverges from the overall distribution (our pangenome) adds nuance to our approach. Summed across all genomes, the relative entropy gives another, complementary angle on pangenome complexity and compression.

To formalise this concept, we turn to bedrock principles in ecology. Ecology has an extensive history of incorporating ideas from information theory and compression (Chao et al. 2013). The Shannon Index has long been used to combine the effects of species richness, the absolute number of unique species in an environment, and species evenness, the relative abundances of those species (Jost 2010). But for ecologists the core equation is more general than Shannon entropy. Ecological datasets span many types of environments. Comparing diversity across those

environments is crucial. Hill introduced the effective number of species as an intuitive solution (Hill 1973). The effective number of species of order q is given as

$$D_q = \left(\sum_{i=1}^N p_i^q \right)^{1/(1-q)}$$

where p_i is the frequency of a particular species i , and N is the total number of unique species. Sweeping through the parameter q controls the metric's responsiveness to rare ($q=0$) or common species ($q=2$ or more). At $q=0$, the expression reduces to species richness, and at $q=2$, the expression expands into the Simpson Index. But the sweet spot is at $q=1$. The limit as q approaches 1 is Shannon's information entropy (or the Shannon Index if overheard in an ecology department). This transformation connects ideas from mathematical ecology to information theory. The exponent of Shannon entropy yields the Hill number at $q=1$, or the effective number of species:

$$D_1 = \exp \left(- \sum_{i=1}^N p_i \ln p_i \right)$$

More diverse samples have higher Hill numbers. Hill numbers convey species diversity as an intuitive metric. Because of the connection to information theory, Hill numbers are not the exclusive domain of ecologists. In Natural Language Processing, perplexity (Brown et al. 1992) is used to measure how well a language model can predict a string of text. Perplexity is the effective number of words in a library. Perplexity and Hill numbers draw from the same mathematical toolkit. This toolkit's simplicity allows for easy comparisons between entirely different experiments. But the expression collapses each experiment's observations into a single distribution.

We can enrich Hill numbers by extending beyond species measured as single variable distributions. To this point, we've defined

what an ecologist would term alpha diversity (Jost 2007), or the diversity of species in any one sample. One sample usually doesn't cut it. Ecologists sample multiple transects from their environment of interest. Sampling introduces several opportunities. First, the degree of sample overlap is a potential gauge of efficacy. Second, sample diversity yields insight into the overall, hypothetical, unapproachable diversity of the system, or the gamma diversity. If samples are highly diverse, ascertaining the diversity of the target environment may require more samples. If gamma diversity is too high, no sampling scheme may be enough. Beta diversity measures the degree of overlap between samples (Marcon et al. 2012; Marcon et al. 2014). Grounding the concept in information theory, we extend the Hill number of species into a Hill number of samples. The following expression yields the effective number of samples:

$$D_{\beta} = \exp \left(\sum_{s=1}^M w_s \sum_{i=1}^N p_{si} \ln \frac{p_{si}}{p_i} \right)$$

This equation incorporates the Kullback–Leibler divergence or relative entropy, a formulation as frequently used as entropy in the information theory literature (Kullback and Leibler 1951). The relative entropy measures the divergence of any one species distribution against the distribution of all species. Here, N is the number of unique species, M , the number of samples, p_{si} the frequency of species i in sample s , p_i the frequency of species i across all samples, and w_s the weight all observations in sample s relative to all individuals collected in the experiment.

The effective number of samples is another measure of compression. If species richness and evenness are the same across all samples, the effective number of samples reduces to 1. If the samples contain no species in common, or if samples have wildly different species occurrence counts, the effective number of samples approaches the number of samples taken. In the first case, we have perfect compression; in the latter, no compression at all.

We take this ecological concept and adapt it to genomics. Our goal is to calculate the information diversity embedded in sequence ensembles. This requires a complete reframe. Rather than species in a community (alpha diversity), we think k-mers in a genome. Rather than transects in an environment (beta diversity), we think genomes in a pangenome. The shift is in the container. Employed in this way, we recast Hill numbers as the effective number of genomes or genome equivalents. We coin KHILL (Narechania et al. 2024), an intuitive metric that quantifies the information space of a pangenome, or the degree to which it will compress. Because KHILL is weighted, we can compare statistics across species regardless of how many genomes are available for each. This allows us to compare the information diversity of pathogens alongside less represented organisms from the rare biosphere. We calculate KHILLs in a fraction of the time it takes to annotate genomes, run alignments and build the orthologues required to compute pangenome fluidity.

5 | The Toolkit Applied

Biological datasets are large and growing. Other fields also contend with large datasets, and some have been grappling with them for decades longer. For example, astronomers have

big data, perhaps the biggest data in the sciences (Stephens et al. 2015). It is impossible to process and save all astronomic data. Astronomers have known for years the importance of sensing data as it shines onto their mirrors. Compression normally happens at the point of collection. We are quickly reaching this point in biology.

Organised, collaborative genome sequencing projects began in earnest in the 1990s. Starting then and through the first two decades of this century, genomic datasets were sacrosanct. Groups held fast to their data until every angle was exhausted. Though genomic data have always been big data, generating it back then was costly. This is no longer the case. The price of genome sequencing has seen a steep decline. Storing these accumulating data has become nearly impossible. Perhaps it is time to let go. With the information bottleneck, we stream data through a channel encouraging controlled data loss. New sequencing platforms emit data in nearly unending streams. Sensors are designed to glean information from data streams in real time. There are sensors that detect change in acceleration (engineers), in light (astronomers), in brain activity (doctors). Perhaps streams of biological sequences can also be processed and discarded. Can a sequence become a sensor?

Take for example SARS-CoV-2. Fifteen million SARS-CoV-2 genomes are now available in various repositories around the world (Shu and McCauley 2017). The state of the art in surveilling these genomes as they accumulate in time and space is phylodynamic (Grenfell et al. 2004; Volz et al. 2013). Phylodynamics is the study of organism spread over short time scales using molecular phylogenetics. But phylodynamics is retrospective. Investigators curate a fraction of the genomes available, compare them against an even more rarefied set of references and embed the new alongside the old either in phylogenetic trees or networks. Alignment is the linchpin in this arrangement. Genome alignments feed tools like Nextstrain (Hadfield et al. 2018), which employ Bayesian and likelihood phylogenetic approaches—some of the most computationally costly algorithms in bioinformatics—to extend our view of SARS-CoV-2 biology slightly beyond the anointed references in a database.

We find this limiting. We can use KHILL to look forward, analysing outright all available sequences (Narechania et al. 2024). Whether it's 15 million clinical genomes or streams of wastewater, KHILL is capable of processing terabytes of streaming sequence and flagging the emergence of new variants without relying on the references that confine biological novelty. KHILL can also achieve rapid community analysis as exemplified in our studies of the microbial shifts in the making of cheese (Rodriguez 2025), and the microbiome perturbations caused by broad spectrum antibiotics (unpublished data). Whether it's a life-threatening virus or the cheese you spread on crackers, we use *all* sequence data, not just the bits that will stick to existing references.

For SARS-CoV-2, we calculate one KHILL number per day along a pandemic time course. Compiling these genome equivalents yields an information diversity curve through time. KHILL increases as variants of concern ascend in a population mixing with a prior background. KHILL decreases once these variants

grow dominant and sweep away all other genomic heterogeneity. In this way, we detect the emergence of concerning strains well before annotation clearinghouses have blessed new database entries.

As a genomic measure of compression, KHILL also naturally lends itself to the analysis of pangenomes. In fact, in pandemic surveillance of SARS-CoV-2, we used KHILL as a rolling measure of pangenome complexity. Because of their contracted timeline, pandemic genomes occupy a small information space. The KHILL of all the millions of sequenced SARS-CoV-2 compresses to about 1.15 effective genomes. But KHILL is not restricted to any one biological scale. We can measure the complexity of strains, species, genus and perhaps collections at even higher taxonomic levels. For example, we have used KHILL to calculate the pangenome complexity of all known bacterial species (Narechania, Bobo, Gilbert, and Gopalakrishnan 2025). An analysis at this scale is impossible with current alignment-based bioinformatic techniques. But because KHILL is fast and simple, we can compute genome equivalents for every species in the RefSeq database (version 223). We couple this with metrics derived from block entropy curves (H_{μ} , EE and TI) to calculate the information space occupied by bacterial species. This information theoretic approach democratises species classification, labelling each pangenome across the microbial tree of life with a single number.

As we have defined it, KHILL species complexity mixes two separate phenomena. First, species definitions vary. The Linnaean taxonomy imposes a hierarchy on life, but this hierarchy is not uniformly applied. Species in one part of the taxonomic tree may not mean the same thing to its experts as species in another part of the tree. This is cultural. But it does influence the relative breadth of species buckets. We expect some variation in KHILL based just on these very human inconsistencies.

More interesting, however, is our second observation. Pangenome fluidity (Kislyuk et al. 2011) has been shown to track with some gross aspects of bacterial phenotype (Dewar et al. 2024). For example, host-bound species accustomed to a uniform environment typically have less complex pangenomes, whereas cosmopolitan species occupying diverse niches tend towards more pangenome diversity. Obligate bacteria are less complex than their facultative counterparts. Non-motile organisms are less complex than those on the move. Complexity in these cases was measured as pangenome fluidity. Pangenome fluidity is as near to measuring information-theoretic complexity as alignment-based techniques can get. We find that KHILL, a more direct, swifter measure of complexity, also corresponds to bacterial lifestyle. For example, pathogens have significantly lower KHILL than mutualists. Challenging environments presumably encourage the accretion of pangenome complexity as species contend with instability. Our compression-based techniques squeeze this information from genomes without the normal bioinformatics playbook. We provide a systems biology measure of complexity that compiles the effects of selective pressure on numerous genes and loci, either streamlining or expanding genomes as required of a microbe's lifestyle.

6 | Challenges? in Sacrifice There is Clarity!

Metrics based in compression can distort mechanisms. KHILL increases with population heterogeneity, as in the case of our SARS-CoV-2 populations. But it also increases with genetic distance. This genetic diversity could be the result of environmental pressure, or it could simply be lazy, inconsistent categorisation. Because block entropy curves and KHILL dispense with alignment, we also lose the ability to pinpoint changes in genomic space. Compression obscures mechanisms, sacrifices genome location information, and conflates biological forces. For example, since programmed ribosomal frameshifting in viruses depends on the reading frame, our method may obscure frame-dependent signals in viral studies. And compression will likely collapse individual genome copy number variation. But in a field saturated with sequence data, our approach allows researchers to skim data streams without resorting to the heaviest, most cumbersome algorithms in bioinformatics.

The idea of conflating signal is a hallmark of information-based approaches. Shannon's communication problem is emblematic of this compromise. Distortion is inevitable as information is relayed from sender to receiver. This concept has been used in everything (Vaseghi 2001) from telecommunications, to thermodynamics, to data encoding in Natural Language Processing. More complex data require a broader channel to communicate. But sometimes we must sacrifice nuance for meaning. In fact, compressing away the noise can sometimes distil signal. In other words, conflation sometimes yields clarity.

We take this concept to phylogenomics (Narechania, Bobo, DeSalle, et al. 2025). Mutation, homologous recombination and horizontal gene transfer all distort phylogenomic signal. We can capture the degree of distortion by measuring how difficult it is to compress strings (k-mers) from a set of genomes through the information bottleneck, and into a set number of clusters. If the compression is easy, we need fewer clusters—a narrower channel—to achieve communication at an acceptable level of distortion. But if the genomes are labile, we need more clusters to communicate the added information diversity. The information bottleneck (Tishby et al. 2000) therefore also quantifies complexity.

The clusters that comprise our information channel, are datasets that sort meaning. Where KHILL is a mark of compression, these clusters are actual compressed representations. We can measure the fidelity of the original 'message' carried by the genomes relative to these compressed representations. Clonal, tree-like, bifurcating species generally require fewer clusters to model modes of genomic change. Recombinogenic species require more clusters to achieve the same signal clarity. Like KHILL, this approach conflates biological phenomena. Lossy compression through the bottleneck does not distinguish between mutation and recombination. But for both KHILL and the bottleneck, the compressibility of a set of genomes becomes a metric that can be used to compare sets of species. We anticipate that in future work, both techniques will operate on raw reads, making assembly as optional as alignment. Building evolutionary models from streamed sequence would realise our ambition to sense change directly from raw data.

We began this essay bemoaning genomics as a retrospective enterprise. We believe information theory allows us to shift our gaze forward. Eliminating references opens us to novelty. Decentering the gene offers a new view of pangenome complexity. And eliminating alignment boosts speed. Together these efficiencies recast sequencing as a sensor delimiting change. We can sense change along a pandemic trajectory. We can predict bacterial lifestyle from compression. And we can probe the unbalanced hierarchies of bacterial taxonomy.

Author Contributions

A.N. conceived the project and drafted the original and final versions; S.G. co-wrote the manuscript; M.T.P.G. co-wrote the manuscript.

Conflicts of Interest

The authors declare no conflicts of interest.

Data Availability Statement

The data that support the findings of this study are available in NCBI at <https://www.ncbi.nlm.nih.gov>. These data were derived from the following resources available in the public domain: - RefSeq, <https://www.ncbi.nlm.nih.gov/refseq/>.

References

- Baker, R. H., and R. DeSalle. 1997. "Multiple Sources of Character Information and the Phylogeny of Hawaiian Drosophilids." *Systematic Biology* 46: 654–673.
- Bentz, C., and D. Alikaniotis. 2016. "The Word Entropy of Natural Languages." *arXiv*. <https://doi.org/10.48550/ARXIV.1606.06996>.
- Bialek, W., I. Nemenman, and N. Tishby. 2001. "Predictability, Complexity, and Learning." *Neural Computation* 13: 2409–2463.
- Bialek, W., F. Rieke, R. R. De Ruyter Van Steveninck, and D. Warland. 1991. "Reading a Neural Code." *Science* 252: 1854–1857.
- Boltzmann, L. 1877. "On the Relation Between the Second Law of the Mechanical Theory of Heat and the Probability Calculus With Respect to the Theorems of Thermal Equilibrium." *Wiener Berichte* 76: 373–435.
- Brown, P., S. D. Pietra, V. D. Pietra, J. Lai, and R. Mercer. 1992. "An Estimate of an Upper Bound for the Entropy of English." *Computational Linguistics* 18, no. 1: 31–40.
- Chao, A., Y. T. Wang, and L. Jost. 2013. "Entropy and the Species Accumulation Curve: A Novel Entropy Estimator via Discovery Rates of New Species." *Methods in Ecology and Evolution* 4: 1091–1100.
- Crutchfield, J. P., and D. P. Feldman. 2003. "Regularities Unseen, Randomness Observed: Levels of Entropy Convergence." *Chaos (Woodbury, N.Y.)* 13: 25–54.
- Dewar, A. E., C. Hao, L. J. Belcher, M. Ghoul, and S. A. West. 2024. "Bacterial Lifestyle Shapes Pangenomes." *Proceedings of the National Academy of Sciences of The United States of America* 121: e2320170121.
- Doolittle, W. F. 1999. "Phylogenetic Classification and the Universal Tree." *Science* 284: 2124–2128.
- Eizenga, J. M., A. M. Novak, J. A. Sibbesen, et al. 2020. "Pangenome Graphs." *Annual Review of Genomics and Human Genetics* 21: 139–162.
- Erlich, Y. 2015. "A Vision for Ubiquitous Sequencing." *Genome Research* 25: 1411–1416.

- Grenfell, B. T., O. G. Pybus, J. R. Gog, et al. 2004. "Unifying the Epidemiological and Evolutionary Dynamics of Pathogens." *Science* 303: 327–332.
- Guigó, R. 2023. "Genome Annotation: From Human Genetics to Biodiversity Genomics." *Cell Genomics* 3: 100375.
- Hadfield, J., C. Megill, S. M. Bell, et al. 2018. "Nextstrain: Real-Time Tracking of Pathogen Evolution." *Bioinformatics* 34: 4121–4123.
- Hill, M. O. 1973. "Diversity and Evenness: A Unifying Notation and Its Consequences." *Ecology* 54: 427–432.
- Jaynes, E. T. 1957. "Information Theory and Statistical Mechanics." *Physics Review* 106: 620–630.
- Jost, L. 2007. "Partitioning Diversity Into Independent Alpha and Beta Components." *Ecology* 88: 2427–2439.
- Jost, L. 2010. "The Relation Between Evenness and Diversity." *Diversity* 2: 207–232.
- Kingma, D. P., and M. Welling. 2013. "Auto-Encoding Variational Bayes." *arXiv*. <https://doi.org/10.48550/ARXIV.1312.6114>.
- Kislyuk, A. O., B. Haegeman, N. H. Bergman, and J. S. Weitz. 2011. "Genomic Fluidity: An Integrative View of Gene Diversity Within Microbial Populations." *BMC Genomics* 12: 32.
- Krizhevsky, A., I. Sutskever, and G. E. Hinton. 2017. "ImageNet Classification With Deep Convolutional Neural Networks." *Communications of the ACM* 60: 84–90.
- Kullback, S., and R. A. Leibler. 1951. "On Information and Sufficiency." *Annals of Mathematical Statistics* 22: 79–86.
- Li, H., and R. Durbin. 2009. "Fast and Accurate Short Read Alignment With Burrows–Wheeler Transform." *Bioinformatics* 25: 1754–1760.
- Li, H., B. Handsaker, A. Wysoker, et al. 2009. "The Sequence Alignment/Map Format and SAMtools." *Bioinformatics* 25: 2078–2079.
- Marcon, E., B. Hérault, C. Baraloto, and G. Lang. 2012. "The Decomposition of Shannon's Entropy and a Confidence Interval for Beta Diversity." *Oikos* 121: 516–522.
- Marcon, E., I. Scotti, B. Hérault, V. Rossi, and G. Lang. 2014. "Generalization of the Partitioning of Shannon Diversity." *PLoS One* 9: e90289.
- Narechania, A., D. Bobo, K. Deitz, R. DeSalle, P. J. Planet, and B. Mathema. 2024. "Rapid SARS-CoV-2 Surveillance Using Clinical, Pooled, or Wastewater Sequence as a Sensor for Population Change." *Genome Research* 34: 1651–1660.
- Narechania, A., D. Bobo, R. DeSalle, B. Mathema, B. Kreiswirth, and P. J. Planet. 2025. "What Do we Gain When Tolerating Loss? The Information Bottleneck Wrings out Recombination." *Molecular Biology and Evolution* 42: msaf029.
- Narechania, A., D. Bobo, M. T. P. Gilbert, and S. Gopalakrishnan. 2025. "The Information Content of Species: Formal Definitions of Pangenome Complexity Track With Bacterial Lifestyle." *bioRxiv*: 2025. <https://doi.org/10.1101/2025.03.28.645969>.
- Neumann, J. S., R. Desalle, A. Narechania, B. Schierwater, and M. Tessler. 2021. "Morphological Characters Can Strongly Influence Early Animal Relationships Inferred From Phylogenomic Data Sets." *Systematic Biology* 70: 360–375.
- Rodriguez, J. A. 2025. "The Effect of Different Milk Pretreatment Methods on Microbiome Community Development During Herrgards Cheese Production and Ripening." *bioRxiv*: 2025. <https://doi.org/10.1101/2025.04.23.648337>.
- Rokas, A., B. L. Williams, N. King, and S. B. Carroll. 2003. "Genome-Scale Approaches to Resolving Incongruence in Molecular Phylogenies." *Nature* 425: 798–804.

- Shannon, C. E. 1948. "A Mathematical Theory of Communication." *Bell System Technical Journal* 27: 379–423.
- Shi, T., and P. G. Falkowski. 2008. "Genome Evolution in Cyanobacteria: The Stable Core and the Variable Shell." *Proceedings of the National Academy of Sciences* 105: 2510–2515.
- Shu, Y., and J. McCauley. 2017. "GISAID: Global Initiative on Sharing All Influenza Data – From Vision to Reality." *Eurosurveillance* 22.
- Stephens, Z. D., S. Y. Lee, F. Faghri, et al. 2015. "Big Data: Astronomical or Genomical?" *PLoS Biology* 13: e1002195.
- Tettelin, H., V. Masignani, M. J. Cieslewicz, et al. 2005. "Genome Analysis of Multiple Pathogenic Isolates of *Streptococcus agalactiae*: Implications for the Microbial "Pan-Genome"." *Proceedings of the National Academy of Sciences* 102: 13950–13955.
- The ENCODE Project Consortium. 2012. "An Integrated Encyclopedia of DNA Elements in the Human Genome." *Nature* 489: 57–74.
- Thomas, A. M., and N. Segata. 2019. "Multiple Levels of the Unknown in Microbiome Research." *BMC Biology* 17: 48.
- Tishby, N., F. C. Pereira, and W. Bialek. 2000. "The Information Bottleneck Method." *arXiv*. <https://doi.org/10.48550/ARXIV.PHYSICS/0004057>.
- Vaseghi, S. V. 2001. *Advanced Digital Signal Processing and Noise Reduction.*, 29–43. John Wiley & Sons. <https://doi.org/10.1002/0470841621.ch2>.
- Volz, E. M., K. Koelle, and T. Bedford. 2013. "Viral Phylodynamics." *PLoS Computational Biology* 9: e1002947.
- Wu, D., P. Hugenholtz, K. Mavromatis, et al. 2009. "A Phylogeny-Driven Genomic Encyclopaedia of Bacteria and Archaea." *Nature* 462: 1056–1060.
- Zhu, Q., U. Mai, W. Pfeiffer, et al. 2019. "Phylogenomics of 10,575 Genomes Reveals Evolutionary Proximity Between Domains Bacteria and Archaea." *Nature Communications* 10: 5477.