

What we talk about when we talk about microbial species

Running title: Compressing genomes to understand their evolution

Apurva Narechania^{1,2}, Shyam Gopalakrishnan², M Thomas P Gilbert^{2,3}

¹Institute for Comparative Genomics, American Museum of Natural History, New York, NY, USA

²Center for Evolutionary Hologenomics, The Globe Institute, University of Copenhagen, Copenhagen, Denmark

³University Museum, NTNU, Trondheim, Norway

Corresponding authors: anarechania@amnh.org; tgilbert@sund.ku.dk

Keywords: Compression, information theory, pangenomes

Glossary

Sequence ensemble

A population-level view of genomic sequences treated as a statistical ensemble, rather than as alignments. This framing enables direct measurement of diversity, structure, and novelty without relying on reference coordinates.

Pangenome

The complete collection of genes and genomic elements across all members of a species or clade, including both the conserved core genome and the variable accessory genome.

Compression

The process of representing data more efficiently by removing redundancy. In comparative genomics, compression highlights shared structure across genomes, enabling reference-free comparisons and novelty detection.

Entropy

A measure of uncertainty or unpredictability. In our context, entropy captures the genomic information content of a population, with higher entropy reflecting greater information diversity.

Relative entropy (Kullback–Leibler divergence)

A measure of the difference between two or more probability distributions. In genome analysis, relative entropy quantifies how much the distribution of sequence features from single genomes deviate from the pangenome.

Information bottleneck

A lossy compression method that reduces data while preserving the information most relevant for predicting a target variable. Applied to genomics, the bottleneck compresses k-mers or other features into coarser clusters while retaining genome origin information. The resulting clusters distill evolutionary signal.

1 *Abstract*

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

21

22 *Introduction (The problem)*

23 Compression encodes information into reduced representations. Whether bits are
24 eliminated through statistical redundancy (lossless compression), or shed entirely (lossy
25 compression), compressed data always has a smaller footprint than the original. The act of
26 compression – its difficulty or ease – communicates information about the original data source.

27 Highly redundant data with many common patterns will compress easily. In contrast, novelty or
28 surprise with little repeated context is difficult to compress. Evolution creates ensembles of
29 sequence. These ensembles can be represented as pangenomes. Pangenomes are compressible
30 entities, but how compressible depends on evolutionary strategy.

31 Genomics is a retrospective field. Existing bioinformatic techniques often model new
32 genomes on sequences annotated in the past¹. Alignment to these reference genomes
33 circumscribes our knowledge of diversity. Large swaths of the tree of life are presumably
34 unknown². For example, much of the sequence from environmental samples passes through
35 annotation filters as undefined³. In a read streaming era⁴, we need forward looking techniques
36 that flag genomic novelty by dispensing with references, annotation or alignment. Standard
37 methods are ill-equipped for these volumes. New species are not easily caught in the sparse web
38 of the known.

39 As genomics has swept through biology, systematics has come to favor molecular
40 character sets to help delimit species boundaries^{5,6}. While morphology is still important, and
41 holdouts have been more than vocal⁷, phylogenomics has more recently carried the day.
42 Phylogenomics extends the handful of marker genes that were the foundation of early molecular
43 systematics to matrices that concatenate thousands of orthologous genes⁸. This character
44 explosion has been a boon to systematics, but gene annotation is still anchored to the known.

45 These orthologous genes are rarely evenly distributed among the genomes that describe a
46 species⁹. The complete set is one definition of the pangenome, and its complexity was originally
47 defined as the rate of gene accumulation with newly sequenced genomes¹⁰. Genes found
48 universally comprise a genomic core and are considered indispensable for basic species
49 functions. Genes found sporadically may contribute to strain success in particular niches but may
50 not be essential to their overall biology. The ratio of core genomes to accessory genomes informs

51 genome fluidity¹¹. Species whose genomes are mostly core have closed, less fluid pangenomes.
52 Species with a large fraction of accessory genes are considered open and more fluid.

53 This gene-centric view of orthologs is blind to the diversity in the non-coding genome¹².
54 Whole genome alignment to annotated, chromosomal references^{13,14} makes variation in non-
55 coding genome accessible but again circumscribes its characterization. If all we know is a linear
56 reference on a single coordinate system, our understanding of the non-coding genome will be
57 limited to what will stick.

58 More recent pangenome methods attempt to enhance the reference by conveying it as a
59 graph¹⁵. For example, a species graph through elements of the genomic core would collapse into
60 a single consensus, punctuated by bubbles that code small scale variation like single nucleotide
61 polymorphism and small insertion/deletion elements. In contrast, the accessory genome forks the
62 pangenome graph along entirely disparate paths. Graph-based methods attempt to incorporate
63 nuance and novelty into a more complex reference structure. But the game is still the same: new
64 data is aligned to a set of old genomes bound together into a complex, branching network.

65 Is there another way? Can we measure some other property of whole genomes that isn't
66 contingent on their alignment? Can we de-center the gene so we aren't limited to the protein
67 coding genome? Can we dispense with phylogenomics so we aren't spending CPU years
68 deciphering a bifurcating set of species relationships that convey a mere shadow of a more
69 reticulate truth¹⁶?

70 Here, we propose several new information theoretic techniques that reimagine genomes
71 as ensembles of information, containers subject to compression. This view of genomic
72 information does not require annotation. Because we aren't concerned with genes or the
73 contiguous arrangements of genomic elements, we also forgo alignment. We instead describe
74 pangenomes with summary statistics of string-based intersections. In this article, we argue that
75 compression can enhance existing comparative genomic strategies, highlight structural evolution

76 through controlled information loss, and democratize the bacterial species question by applying a
77 uniform mathematics across the Linnean taxonomy.

78

79 *The toolkit (entropy)*

80 Our approach is guided by two foundational concepts at the very root of information
81 theory: entropy and relative entropy. Both ideas rely heavily on Claude Shannon's seminal ideas
82 on information introduced in "A Mathematical Theory of Communication", the founding
83 document of information theory¹⁷. Information is data that reduces uncertainty. Shannon's
84 original formulation resembled the thermodynamic construction of entropy devised for statistical
85 mechanics¹⁸. We measure information entropy as

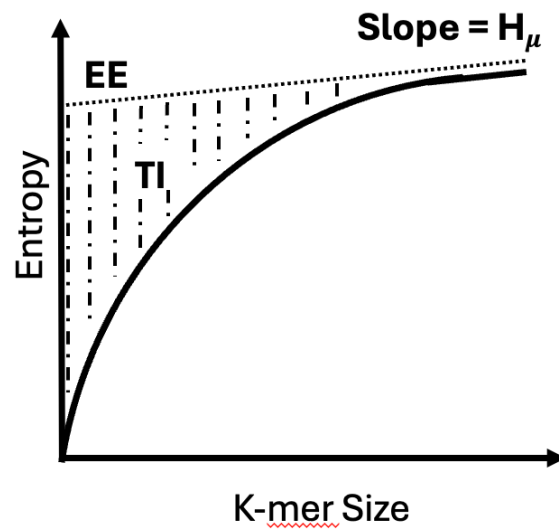
86
$$H = - \sum_{i=1}^N p_i \ln p_i$$

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

90 In genomics, our data comes in sequences. We can measure the entropy of sequences by
91 digesting into substrings of specific size. In the bioinformatics literature, substrings of biological
92 readouts (DNA, RNA, protein) are called k-mers. In a comparative setting, we're most interested
93 in the entropy of a group of sequences, or sequence ensemble^{19,20}. For genome sequence
94 ensembles, alignment has been the tool of choice. But alignment is computationally arduous and
95 breaks down with evolutionary distance. Fields as diverse as linguistics²¹, neurobiology²², and
96 statistical mechanics²³ have successfully employed entropy to quantify ensemble complexity. In
97 each of these fields, researchers code a linear string of observations and divide into
98 subsequences, calculating the entropy of each set across the ensemble. In Figure 1, we show how

99 the entropy of genome sequence (e.g. DNA/RNA) typically increases with increasing
100 subsequence size. This is a block entropy curve.

101 Block entropy curves contain information about the complexity of the ensemble¹⁹.
102 Systems with more ensemble structure – repeated elements across sequences – will peak at lower
103 entropy. More novelty across sequences yields higher entropy. In genomics, closed pangenomes,
104 with a large core shared across all species
105 genomes, have low entropy. Auxiliary genes
106 unique to subsets of genomes add entropy to
107 the ensemble system. The uneven distribution
108 of these elements is the hallmark of an open
109 pangenome. But to measure complexity we
110 don't need the annotated and aligned genes.



111 Signal is preserved in unaligned and
112 unannotated k-mers.

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).

113 Block entropy curves asymptote at the
114 minimum block size required to efficiently capture information diversity across the sequences.
115 We use three quantities calculated from these curves to describe the complexity of a pangenome:
116 source entropy, excess entropy and transient information¹⁹. The source entropy (H_{μ}) is the
117 irreducible randomness that remains even as larger block sizes capture most ensemble
118 correlations. H_{μ} is a direct measure of randomness. Random distributions are hard to
119 compress. A high source entropy is associated with the accumulation of unevenly distributed
120 accessory genes, resulting in a more complex pangenome. The excess entropy (EE) is the non-
121 random fraction of the total information in the system. It's the information we model from
122 redundancies across the ensemble. Alignment is anchored to these redundancies. In fact,
123 alignment only works if enough of these redundancies are spread across the query genomes.

124 Finally, the transient information (TI) measures how much information we must invest to learn
125 Hmu and EE. In Figure 1 we show TI as the area between the block entropy curve and the line
126 defining Hmu. Species with closed pangenomes typically have a lower TI than those open to
127 accumulating gene diversity. Closed pangenomes with a large core set of genes compress at
128 lower k-mer sizes, approaching their Hmu quickly.

129

130 *More tools in the toolkit (the information bottleneck)*

131 Entropy is the workhorse of lossless compression. In fact, it defines lossless
132 compression's limit. We cannot compress any further than the entropy of the source. In our
133 context, the block entropy curve follows compression limits along a k-mer spectrum. Lossless
134 compression preserves all data, but sacrifice can feature evolutionary events by isolating patterns
135 from genomic noise. Using lossy compression, we can identify the core genome of any species
136 without alignment or annotation. Along the way, we unlock the homologous and non-
137 homologous recombination events that violate vertical signal⁴⁴.

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

144 No channel is noise-less. Still, the distortion introduced by noisy channels does not doom
145 message passing. A sender can compensate for noise by encoding more information into a
146 message, or a receiver can tolerate some level of distortion while ascertaining a sender's core
147 meaning. Shannon formalized this concept as rate-distortion theory¹⁷. On the sender side, the rate
148 is measured as bits of information per symbol. The sender's message is distorted as it passes

149 through the channel. The sender's rate and the receiver's distortion are inversely related. The
 150 function describing the two variables for any given channel informs lossy compression. How
 151 much information loss can we tolerate in reconstructing a sender's message?

152 This idea is central not only to information theory and lossy coding, but also to modern
 153 machine learning methods that use variational autoencoders to populate the compressive layers
 154 of a neural net^{24,25}. The two key questions are 1. How well does a dataset compress, and 2. How
 155 much data can we afford to lose?

156 We use these concepts to further understand pangenome complexity. Imagine the
 157 compression regime in Figure 2. A set of genomes comprising a sequence ensemble are digested
 158 into k-mers and compressed into a set number of clusters. This compression is analogous to a
 159 communication channel. The more clusters we model, the higher the rate, and the lower the

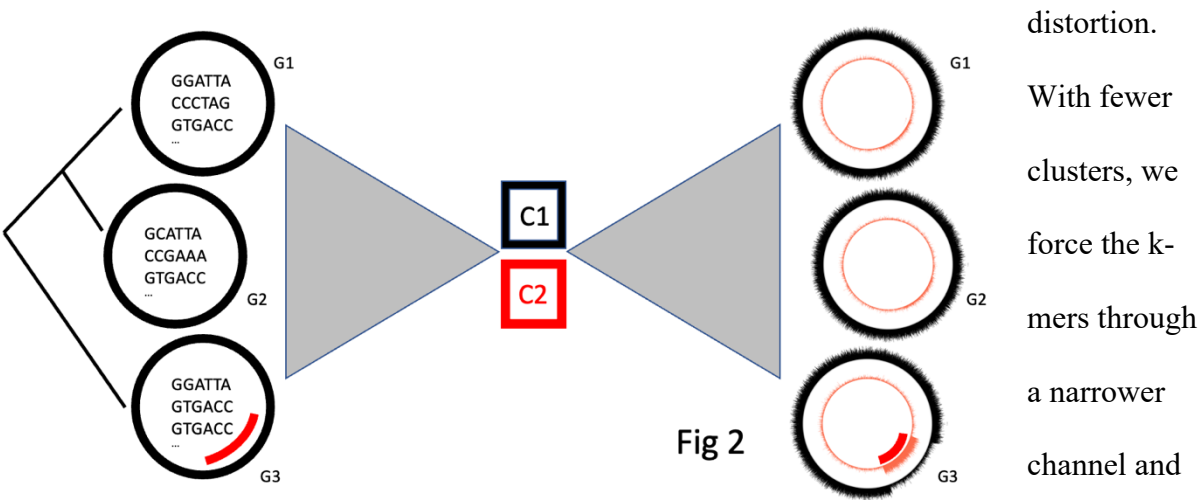


Figure 2. The information bottleneck. As k-mers from our input genomes are compressed into a narrow channel, patterns of structural evolution emerge from the resulting clusters.

169 If we hold the channel constant and model the same number of clusters across species
 170 ensembles, open pangenomes will suffer more data loss than their closed counterparts. Open
 171 pangenomes have more complex information to communicate. This is the information
 172 bottleneck²⁶, an idea first proposed in the Natural Language Processing literature. The bottleneck
 173 modulates loss in our compression framework. Moreover, the clusters we glean comprise a

174 model of structural evolution. The largest cluster usually represents the core genome. K-mers
175 from recombination regions populate the others. For example, we have shown that we can detect
176 the horizontal gene transfer events that fueled the evolution of pathogenicity in *Staphylococcus*
177 *aureus* Clonal Complex 8 (CC8). We do this *de novo*, without knowing anything about the
178 elements themselves. In CC8, the staphylococcal chromosomal cassette carrying the methicillin
179 resistance gene expanded, accumulating other disease associated elements as the pathogen
180 moved between North and South America⁴⁴. The act of compression therefore learns real
181 biological events without the need to annotate genes, align genomes, build sets of orthologs, or
182 calculate any trees.

183

184 *Even more tools in the toolkit (relative entropy)*

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

195 To formalize this concept, we turn to bedrock principals in ecology. Ecology has an
196 extensive history of incorporating ideas from information theory and compression²⁷. The
197 Shannon Index has long been used to combine the effects of species richness, the absolute
198 number of unique species in an environment, and species evenness, the relative abundances of

199 those species²⁸. But for ecologists the core equation is more general than Shannon entropy.
200 Ecological datasets span many types of environments. Comparing diversity across those
201 environments is crucial. Hill introduced the effective number of species as an intuitive solution²⁹.
202 The effective number of species of order q is given as

$$203 \quad D_q = \left(\sum_{i=1}^N p_i^q \right)^{1/(1-q)}$$

204 where p_i is the frequency of a particular species i , and N , the total number of unique species.
205 Sweeping through the parameter q controls the metric's responsiveness to rare ($q = 0$) or
206 common species ($q = 2$ or more). At $q = 0$, the expression reduces to species richness, and at $q =$
207 2 , the expression expands into the Simpson Index. But the sweet spot is at $q = 1$. The limit of this
208 equation as q approaches 1 is Shannon's information entropy (or the Shannon Index if overheard
209 in an ecology department). This transformation connects ideas from mathematical ecology to
210 information theory. The exponent of Shannon entropy yields the Hill number at $q = 1$, or the
211 effective number of species:

$$212 \quad D_1 = \exp \left(- \sum_{i=1}^N p_i \ln p_i \right)$$

213 More diverse samples have higher Hill numbers. Hill numbers convey species diversity as an
214 intuitive number. Because of its connection to information theory, Hill numbers are not the
215 exclusive domain of ecologists. In Natural Language Processing, perplexity³⁰ is used to measure
216 how well a language model can predict a string of text. Perplexity is the effective number of
217 words in a library. Perplexity and Hill numbers draw from the same mathematical toolkit. This
218 toolkit's simplicity allows for easy comparisons between entirely different experiments. But the
219 expression collapses each experiment's observations into a single distribution.

220 We can enrich Hill numbers by extending beyond species measured as single variable
221 distributions. To this point, we've defined what an ecologist would term alpha diversity³¹, or the

222 diversity of species in any one sample. One sample usually doesn't cut it. Ecologists sample
223 multiple transects from their environment of interest. Sampling introduces several opportunities.
224 First, the degree of sample overlap is a potential gauge of efficacy. Second, sample diversity
225 yields insight into the overall, hypothetical, unapproachable diversity of the system, or the
226 gamma diversity. If samples are highly diverse, ascertaining the diversity of the target
227 environment may require more samples. If gamma diversity is too high, no sampling scheme
228 may be enough. Beta diversity measures the degree of overlap between samples^{32,33}. Grounding
229 the concept in information theory, we extend the Hill number of species into a Hill number of
230 samples. The following expression yields the effective number of samples:

231
$$D_{\beta} = \exp \left(\sum_{s=1}^M w_s \sum_{i=1}^N p_{si} \ln \frac{p_{si}}{p_i} \right)$$

232 This equation incorporates the Kullback-Leibler divergence or relative entropy, a formulation as
233 frequently used as entropy in the information theory literature³⁴. The relative entropy measures
234 the divergence of any one genome's k-mer distribution against the k-mer distribution of the
235 entire pangenome. Here, N is the number of unique species, M , the number of samples, p_{si} the
236 frequency of species i in sample s , p_i the frequency of species i across all samples, and w_s the
237 weight all observations in sample s relative to all individuals collected in the experiment.

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

243 We take this ecological concept and adapt it to genomics. Our goal is to calculate the
244 information diversity embedded in sequence ensembles. This requires a complete reframe.
245 Rather than species in a community (alpha diversity), we think k-mers in a genome. Rather than

246 transects in an environment (beta diversity), we think genomes in a pangenome. The shift is in
247 the container. Employed in this way, we recast Hill numbers as the effective number of genomes
248 or genome equivalents. We coin KHILL⁴⁰, an intuitive metric that quantifies the information
249 space of a pangenome, or the degree to which it will compress. Because KHILL is weighted, we
250 can compare statistics across species regardless of how many genomes are available for each.
251 This allows us to compare the information diversity of pathogens alongside less represented
252 organisms from the rare biosphere. We calculate KHILLs in a fraction of the time it takes to
253 annotate genomes, run alignments, and build the orthologs required to compute pangenome
254 fluidity.

255

256 *The toolkit applied*

257 Biological datasets are large and growing. Other fields also contend with large datasets,
258 and some have been grappling with them for decades longer. For example, astronomers have big
259 data, perhaps the biggest data in the sciences³⁵. Processing and saving all astronomic data is
260 impossible. Astronomers have known for years the importance of sensing data as it shines onto
261 their mirrors. Compression normally happens at the point of collection. We are quickly reaching
262 this point in biology.

263 Organized, collaborative genome sequencing projects began in earnest in the 1990s.
264 Starting then and through the first two decades of this century, genomic datasets were sacrosanct.
265 Groups held fast to their data until every angle was exhausted. Though genomic data has always
266 been big data, generating it back then was costly. This is no longer the case. The price of genome
267 sequencing has seen steep decline. Storing this accumulating data has become nearly impossible.
268 Perhaps it is time to let go. With the information bottleneck, we stream data through a channel,
269 and encourage controlled data loss. New sequencing platforms emit data in nearly unending
270 streams. Sensors are designed to glean information from data streams in real time. There are

271 sensors that detect change in acceleration (engineers), in light (astronomers), in brain activity
272 (doctors). Perhaps streams of biological sequence can also be processed and discarded. Can
273 sequence become a sensor?

274 Take for example SARS-CoV-2. Fifteen million SARS-CoV-2 genomes are now
275 available in various repositories around the world³⁶. The state of the art in surveilling these
276 genomes as they accumulate in time and space is phylodynamic^{37,38}. Phylodynamics is the study
277 of organism spread over short time scales with molecular phylogenetics. But phylodynamics is
278 retrospective. Investigators curate a fraction of the genomes available, compare them against an
279 even more rarefied set of references, and embed the new alongside the old either in phylogenetic
280 trees or networks. Alignment is the linchpin in this arrangement. Genome alignments feed tools
281 like Nextstrain³⁹, which employ Bayesian and likelihood phylogenetic approaches – some of the
282 most computationally costly algorithms in bioinformatics – to extend our view of SARS-CoV-2
283 biology slightly beyond the anointed references in a database.

284 We find this limiting. We can use KHILL to look forward, analyzing all the sequence
285 available to us outright⁴⁰. Whether it's 15 million clinical genomes or streams of wastewater,
286 KHILL is capable of processing terabytes of streaming sequence and flagging the emergence of
287 new variants without relying on the references that confine biological novelty. KHILL can also
288 achieve rapid community analysis as exemplified in our study of the microbial shifts in the
289 making of cheese (ref), and the microbiome perturbations caused by broad spectrum antibiotics
290 (unpublished data). Whether it's a life-threatening virus or the cheese you spread on crackers, we
291 use *all* sequence, not just the bits that will stick to existing references.

292 For SARS-CoV-2, we calculate one KHILL number per day along a pandemic time
293 course. Compiling these genome equivalents yields an information diversity curve through time.
294 KHILL increases as variants of concern ascend in a population mixing with a prior background.
295 KHILL decreases once these variants grow dominant and sweep away all other genomic

296 heterogeneity. In this way, we detect the emergence of concerning strains well before annotation
297 clearinghouses have blessed new database entries.

298 As a genomic measure of compression, KHILL also naturally lends itself to the analysis
299 of pangenomes. In fact, with SARS-CoV-2, we used KHILL as a rolling measure of pangenome
300 complexity. Because of their contracted timeline, pandemic genomes occupy a small information
301 space. The KHILL of all the millions of sequenced SARS-CoV-2 compresses to about 1.15
302 effective genomes. But KHILL is not restricted to any one biological scale. We can measure the
303 complexity of strains, species, genus, and perhaps collections at even higher taxonomic levels.

304 For example, we have used KHILL to calculate the pangenome complexity of all known
305 bacterial species⁴¹. An analysis at this scale is impossible with current alignment-based
306 bioinformatic techniques. But because KHILL is fast, we compute genome equivalents for every
307 species in the RefSeq database (version 223). We couple this with metrics derived from block
308 entropy curves (Hmu, EE and TI) to calculate the information space occupied by bacterial
309 species. This information theoretic approach democratizes species classification, labelling each
310 pangenome across the microbial tree of life with a single number.

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

317 More interesting, however, is our second observation. Pangenome fluidity¹¹ has been
318 shown to track with some gross aspects of bacterial phenotype⁴². For example, host-bound
319 species accustomed to a uniform environment typically have less complex pangenomes.
320 Cosmopolitan species occupying diverse niches tend towards more pangenome diversity.

321 Obligate bacteria are less complex than their facultative counterparts. Non-motile organisms,
322 less complex than those on the move. Complexity, in this case, was measured as pangenome
323 fluidity. Pangenome fluidity is as near to measuring information-theoretic complexity as
324 alignment-based techniques can get. We find that KHILL, a more direct, swifter measure of
325 complexity, also corresponds to bacterial lifestyle. For example, pathogens have significantly
326 lower KHILL than mutualists. Challenging environments presumably encourage the accretion of
327 pangenome complexity as species contend with instability. Our compression-based techniques
328 squeeze this information from genomes without the normal bioinformatics playbook. We provide
329 a systems biology measure of complexity that compiles the effects of selective pressure on
330 numerous genes and loci, either streamlining or expanding genomes as required of a microbe's
331 lifestyle.

332

333 *Challenges? In sacrifice there is clarity!*

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

345 The idea of conflating signal is a hallmark of information-based approaches. Shannon's
346 communication problem is emblematic of this compromise. Distortion is inevitable as
347 information is relayed from sender to receiver. This concept has been used in everything⁴³ from
348 telecommunications, to thermodynamics, to data encoding in Natural Language Processing.
349 More complex data requires a broader channel to communicate. But sometimes we must
350 sacrifice nuance for meaning. In fact, compressing away the noise can sometimes distill signal.
351 In other words, conflation sometimes yields clarity.

352 We take this concept to genomics⁴⁴. Mutation, homologous recombination, and
353 horizontal gene transfer all distort genomic signal. We can capture the degree of distortion by
354 measuring how difficult it is to compress strings (k-mers) from a set of genomes, through the
355 information bottleneck, and into a set number of clusters. If the compression is easy, we need
356 fewer clusters – a narrower channel – to achieve communication at an acceptable level of
357 distortion. But if the genomes are labile, we need more clusters to communicate the added
358 information diversity. The information bottleneck²⁶ therefore also quantifies complexity.

359 The clusters that comprise our information channel, are datasets that sort meaning. Where
360 KHILL is a mark of compression, these clusters are actual compressed representations. We can
361 measure the fidelity of the original 'message' carried by the genomes relative to these
362 compressed representations. Clonal, tree-like, bifurcating species generally require fewer clusters
363 to model modes of genomic change. Recombinogenic species require more clusters to achieve
364 the same signal clarity.

365 Like KHILL, this approach conflates biological phenomena. Lossy compression through
366 the bottleneck does not distinguish between mutation and recombination. But for both KHILL
367 and the bottleneck, the compressibility of a set of genomes becomes a metric that can be used to
368 compare sets of species. We anticipate that in future work, both techniques will operate on raw

369 reads, making assembly as optional as alignment. Building evolutionary models from streamed
370 sequence would realize our ambition to sense change directly from raw data.

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

377

378 **Data Accessibility**

379 All data discussed in this paper is freely available online at NCBI
380 (<https://www.ncbi.nlm.nih.gov>).

381

382 **Author Contributions**

383 AN conceived the project and drafted the original and final versions; SG co-wrote the
384 manuscript; MTPG co-wrote the manuscript.

385

386 **References**

- 387 1. Guigó, R. Genome annotation: From human genetics to biodiversity genomics. *Cell*
388 *Genomics* **3**, 100375 (2023).
- 389 2. Wu, D. *et al.* A phylogeny-driven genomic encyclopaedia of Bacteria and Archaea. *Nature*
390 **462**, 1056–1060 (2009).
- 391 3. Thomas, A. M. & Segata, N. Multiple levels of the unknown in microbiome research. *BMC*
392 *Biol.* **17**, (2019).
- 393 4. Erlich, Y. A vision for ubiquitous sequencing. *Genome Res.* **25**, 1411–1416 (2015).

- 394 5. Baker, R. H. & DeSalle, R. Multiple Sources of Character Information and the Phylogeny of
395 Hawaiian Drosophilids. *Syst. Biol.* **46**, 654–673 (1997).
- 396 6. Rokas, A., Williams, B. L., King, N. & Carroll, S. B. Genome-scale approaches to resolving
397 incongruence in molecular phylogenies. *Nature* **425**, 798–804 (2003).
- 398 7. Neumann, J. S., Desalle, R., Narechania, A., Schierwater, B. & Tessler, M. Morphological
399 Characters Can Strongly Influence Early Animal Relationships Inferred from Phylogenomic
400 Data Sets. *Syst. Biol.* **70**, 360–375 (2021).
- 401 8. Zhu, Q. *et al.* Phylogenomics of 10,575 genomes reveals evolutionary proximity between
402 domains Bacteria and Archaea. *Nat. Commun.* **10**, (2019).
- 403 9. Shi, T. & Falkowski, P. G. Genome evolution in cyanobacteria: The stable core and the
404 variable shell. *Proc. Natl. Acad. Sci.* **105**, 2510–2515 (2008).
- 405 10. Tettelin, H. *et al.* Genome analysis of multiple pathogenic isolates of *Streptococcus*
406 *agalactiae*: Implications for the microbial “pan-genome”. *Proc. Natl. Acad. Sci.* **102**, 13950–
407 13955 (2005).
- 408 11. Kislyuk, A. O., Haegeman, B., Bergman, N. H. & Weitz, J. S. Genomic fluidity: an
409 integrative view of gene diversity within microbial populations. *BMC Genomics* **12**, 32
410 (2011).
- 411 12. The ENCODE Project Consortium. An integrated encyclopedia of DNA elements in the
412 human genome. *Nature* **489**, 57–74 (2012).
- 413 13. Li, H. & Durbin, R. Fast and accurate short read alignment with Burrows–Wheeler
414 transform. *Bioinformatics* **25**, 1754–1760 (2009).
- 415 14. Li, H. *et al.* The Sequence Alignment/Map format and SAMtools. *Bioinformatics* **25**, 2078–
416 2079 (2009).
- 417 15. Eizenga, J. M. *et al.* Pangenome Graphs. *Annu. Rev. Genomics Hum. Genet.* **21**, 139–162
418 (2020).

16. Doolittle, W. F. Phylogenetic Classification and the Universal Tree. *Science* **284**, 2124–2128 (1999).

17. Shannon, C. E. A Mathematical Theory of Communication. *Bell Syst. Tech. J.* **27**, 379–423 (1948).

18. Boltzmann, L. 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. *Wien. Berichte* **76**, 373–435 (1877).

19. Crutchfield, J. P. & Feldman, D. P. Regularities unseen, randomness observed: Levels of entropy convergence. *Chaos Interdiscip. J. Nonlinear Sci.* **13**, 25–54 (2003).

20. Bialek, W., Nemenman, I. & Tishby, N. Predictability, Complexity, and Learning. *Neural Comput.* **13**, 2409–2463 (2001).

21. Bentz, C. & Alikaniotis, D. The word entropy of natural languages. Preprint at <https://doi.org/10.48550/ARXIV.1606.06996> (2016).

22. Bialek, W., Rieke, F., De Ruyter Van Steveninck, R. R. & Warland, D. Reading a Neural Code. *Science* **252**, 1854–1857 (1991).

23. Jaynes, E. T. Information Theory and Statistical Mechanics. *Phys. Rev.* **106**, 620–630 (1957).

24. Kingma, D. P. & Welling, M. Auto-Encoding Variational Bayes. Preprint at <https://doi.org/10.48550/ARXIV.1312.6114> (2013).

25. Krizhevsky, A., Sutskever, I. & Hinton, G. E. ImageNet classification with deep convolutional neural networks. *Commun. ACM* **60**, 84–90 (2017).

26. Tishby, N., Pereira, F. C. & Bialek, W. The information bottleneck method. (2000) [doi:10.48550/ARXIV.PHYSICS/0004057](https://doi.org/10.48550/ARXIV.PHYSICS/0004057).

- 442 27. Chao, A., Wang, Y. T. & Jost, L. Entropy and the species accumulation curve: a novel
443 entropy estimator via discovery rates of new species. *Methods Ecol. Evol.* **4**, 1091–1100
444 (2013).
- 445 28. Jost, L. The Relation between Evenness and Diversity. *Diversity* **2**, 207–232 (2010).
- 446 29. Hill, M. O. Diversity and Evenness: A Unifying Notation and Its Consequences. *Ecology* **54**,
447 427–432 (1973).
- 448 30. Brown, P., Pietra, S. D., Pietra, V. D., Lai, J. & Mercer, R. An Estimate of an Upper Bound
449 for the Entropy of English. *CL* (1992).
- 450 31. Jost, L. PARTITIONING DIVERSITY INTO INDEPENDENT ALPHA AND BETA
451 COMPONENTS. *Ecology* **88**, 2427–2439 (2007).
- 452 32. Marcon, E., Hérault, B., Baraloto, C. & Lang, G. The decomposition of Shannon’s entropy
453 and a confidence interval for beta diversity. *Oikos* **121**, 516–522 (2012).
- 454 33. Marcon, E., Scotti, I., Hérault, B., Rossi, V. & Lang, G. Generalization of the Partitioning of
455 Shannon Diversity. *PLoS ONE* **9**, e90289 (2014).
- 456 34. Kullback, S. & Leibler, R. A. On Information and Sufficiency. *Ann. Math. Stat.* **22**, 79–86
457 (1951).
- 458 35. Stephens, Z. D. *et al.* Big Data: Astronomical or Genomical? *PLOS Biol.* **13**, e1002195
459 (2015).
- 460 36. Shu, Y. & McCauley, J. GISAID: Global initiative on sharing all influenza data – from
461 vision to reality. *Eurosurveillance* **22**, (2017).
- 462 37. Grenfell, B. T. *et al.* Unifying the Epidemiological and Evolutionary Dynamics of
463 Pathogens. *Science* **303**, 327–332 (2004).
- 464 38. Volz, E. M., Koelle, K. & Bedford, T. Viral Phylodynamics. *PLoS Comput. Biol.* **9**,
465 e1002947 (2013).

466 39. Hadfield, J. *et al.* Nextstrain: real-time tracking of pathogen evolution. *Bioinformatics* **34**,
467 4121–4123 (2018).

468 40. Narechania, A. *et al.* Rapid SARS-CoV-2 surveillance using clinical, pooled, or wastewater
469 sequence as a sensor for population change. *Genome Res.* **34**, 1651–1660 (2024).

470 41. Narechania, A., Bobo, D., Gilbert, M. T. P. & Gopalkrishnan, S. The information content of
471 species: formal definitions of pangenome complexity track with bacterial lifestyle. Preprint
472 at <https://doi.org/10.1101/2025.03.28.645969> (2025).

473 42. Dewar, A. E., Hao, C., Belcher, L. J., Ghoul, M. & West, S. A. Bacterial lifestyle shapes
474 pangenomes. *Proc. Natl. Acad. Sci.* **121**, e2320170121 (2024).

475 43. Noise and Distortion. in *Advanced Digital Signal Processing and Noise Reduction* 29–43
476 (Wiley, 2001). doi:10.1002/0470841621.ch2.

477 44. Narechania, A. *et al.* What Do We Gain When Tolerating Loss? The Information Bottleneck
478 Wrings Out Recombination. *Mol. Biol. Evol.* **42**, msaf029 (2025).

479