

BACBENCH: EVALUATING GENOMIC LANGUAGE MODELS FOR BACTERIA

Anonymous authors

Paper under double-blind review

ABSTRACT

Bacteria underpin key processes in health, ecology, and biotechnology, yet machine learning in bacterial genomics lacks systematic, large-scale evaluation resources. Current resources are typically limited to single-species datasets, where the small number of available genomes leaves species-specific models underpowered, underscoring the need for approaches that can generalize across the bacterial tree of life. To address this gap, we present BacBench, the first comprehensive benchmark for bacterial genomics. BacBench consists of 11 datasets across 6 tasks, including a newly generated dataset for operon identification derived from long-read RNA sequencing. BacBench covers gene-, system-, and genome-scale prediction tasks, spanning 67k genomes, 17.6k species and 255M proteins. We analyze the performance of state-of-the-art DNA LMs, protein LMs and bacterial LMs and find that while each approach excels at different scales—the existing models fail to accurately predict the bacterial phenotype at a whole-genome level, hampering the translation to high-impact applications such as antibiotic-resistance and bioproduction. Therefore, highlighting the need to develop methods that reason over the context of the entire genomes, exploiting genomic synteny and transfer across species. We outline the key requirements for such models and release a standardized library for preprocessing, embedding, and evaluation, fostering the development of methods that accurately represent bacterial genomes, and enabling reproducible comparison of diverse approaches under a unified framework. By providing the first comprehensive benchmark dedicated to bacterial genomics, BacBench lays the ground-work for developing machine learning models that truly exploit shared evolutionary patterns and generalize across the bacterial tree of life.

1 INTRODUCTION

Bacteria drive indispensable processes in medicine, ecology, and biotechnology (de Steenhuisen Piters et al., 2015; Luo et al., 2024). They produce industrial enzymes and antibiotics (Ariaeenejad et al., 2024; Santos-Júnior et al., 2024), recycle nutrients, and are being engineered for carbon capture and waste remediation (Xu & Jiang, 2024). Unlocking this potential hinges on interpreting bacterial genomes at scale. A machine learning (ML) system that can embed and reason over entire bacterial genomes could predict clinically relevant traits, surface novel enzymes for biomanufacturing, and reveal how genetic variation translates into functional capabilities across species.

Traditionally, ML approaches in bacterial genomics have been species-specific. For example, genome-wide association studies (GWAS) have been successful in identifying genotype–phenotype associations within species, such as antimicrobial resistance or virulence traits (Lees et al., 2016; Power et al., 2017; San et al., 2020). However, species-specific datasets typically include only a few genomes, leaving models statistically underpowered and prone to overfitting given the vast mutation space of bacterial genomes.

Meaningful progress therefore requires models that can share information across species. Such transfer is biologically plausible: every bacterium carries a small core of universal single-copy proteins, and many additional gene families are conserved far beyond the species level (Wang et al., 2022; Lang et al., 2013; Coleman et al., 2021). Leveraging these shared signals allows models to capture the full extent of bacterial diversity, leading to predictions that are both robust and generalizable. Training on genomes from many species, laboratories, and environmental contexts

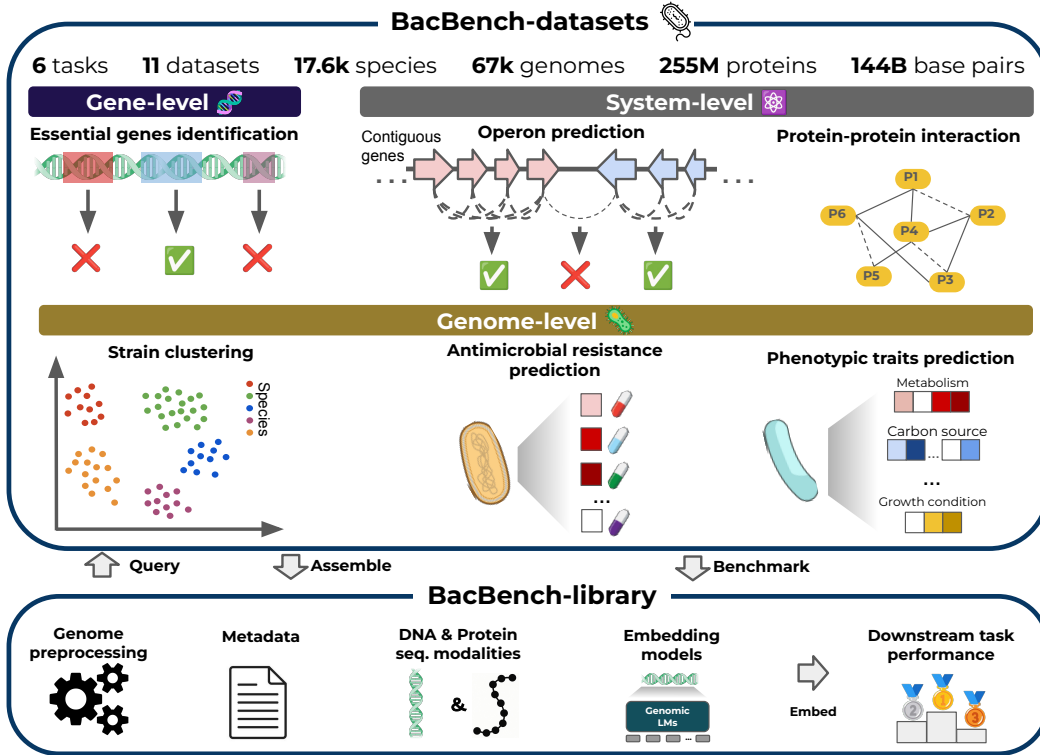


Figure 1: **BacBench overview.** We collected a diverse set of tasks at gene-, system- and genome-scales, spanning 17.6k bacterial species. **BacBench-library** provides support to preprocess and embed the datasets with various models. Finally, we performed systematic benchmarking for each task using diverse genomic LMs.

further guards against dataset-specific artifacts, making the resulting models less prone to sampling bias and more reliable when deployed on previously unseen strains or sequencing pipelines.

Recent breakthroughs in genomic sequence modeling show that the resulting representations can generalize across species and capture evolutionary signals (Nguyen et al., 2024; Dalla-Torre et al., 2024; Zhou et al., 2023; Lin et al., 2022; Elnaggar et al., 2021; Lin et al., 2023; Hayes et al., 2025). However, within the bacterial domain specifically, evaluation has been fragmented—either on narrow, single-task applications like antimicrobial resistance prediction (Wiatrak et al., 2024) or as part of cross-kingdom evaluations that fail to address bacteria-specific challenges (Nguyen et al., 2024). Consequently, the field lacks a dedicated, multi-scale benchmark for bacterial genomics, where genomic and metabolic mechanisms differ substantially from eukaryotes. Thus, leading to the development of models which can accurately model whole bacterial genomes.

Here, we introduce **BacBench**, the first multi-scale, multi-task and multi-species benchmark designed to evaluate ML for bacterial genomics (Fig. 1). BacBench has been collated from a diverse set of 11 datasets organized into 6 tasks spanning multiple biological scales: gene , system  and genome . We consider essential gene prediction task at the gene-level, operon and protein-protein interaction prediction tasks at the system-level; and strain clustering, antibiotic resistance, and phenotypic traits prediction tasks at the genome-level. Collected from a diverse set of public resources and newly generated data, these datasets encompass 67k genomes spanning more than 17.6k bacterial species and 255M proteins. Using BacBench, we conduct comprehensive evaluation of distinct approaches to modeling bacterial genomes including DNA LMs, protein LMs (pLMs) and bacterial LMs (bLMs). We find that different modeling approaches excel at different tasks and scales, yet all models achieve low performance on phenotype prediction at a whole-genome level. Our results demonstrate (i) the need to develop ML approaches which can accurately model entire bacterial genomes, (ii) the benefits of pretraining on bacteria-specific corpora, rather than cross-kingdom ones, thus learning genomic mechanisms that are unique to bacteria, and (iii) the importance of selecting the right model for the task at hand.

BacBench datasets and an accompanying toolkit for preprocessing and evaluation ensure straightforward reuse and extension¹. By generating and integrating these datasets, we introduce the first benchmark suite for bacterial genomics, aiming to catalyze the development of ML methods that can transfer knowledge across species, unlocking new discoveries in bacterial genomics.

2 RELATED WORK

Overall, existing resources for evaluating ML in bacterial genomics are limited in several key ways: they often lack functional labels, restrict evaluation to single species, or evaluate only one biological scale. BacBench addresses these limitations by providing a multi-scale and multi-task benchmark across the bacterial tree of life. It is specifically designed to reflect the core challenges in the field, including data sparsity within individual species, the need for cross-species generalization, and the importance of assessing model performance from the level of individual genes to entire genomes.

Bacterial genomics datasets. Large-scale sequence repositories such as MGnify (Mitchell et al., 2023), IMG/M (Markowitz et al., 2012), and AllTheBacteria (Blackwell et al., 2024) catalogue millions of bacterial assemblies. These resources are indispensable for comparative genomics, yet they provide limited metadata labels, making them ill-suited for training or benchmarking ML models. Simultaneously, task-specific collections provide information on essential genes and phenotypic traits (Zhang et al., 2004; Madin et al., 2020; Brbić et al., 2016; Weimann et al., 2016; Consortium, 2022)—supply richer annotations but usually cover only a handful of species and a single prediction setting. Diverse Genomic Embedding Benchmark (DGEB) (West-Roberts et al., 2024) offers tasks drawn from all domains of life, yet its bacterial coverage is mostly limited to single-species gene or short-segment datasets and lacks genome-scale evaluations such as broad phenotypic traits inference or strain-level clustering, leaving it unable to assess whether models generalize across thousands of species at multiple scales. BacBench complements these efforts by integrating six heterogeneous tasks, ranging from gene essentiality to genome-wide phenotype prediction - into a unified framework that explicitly tests generalization across 17.6 k bacterial species and three biological scales (Fig. 1).

Single-species bacterial genomics models. Despite the existence of large-scale bacterial genomics datasets, the ML applications in bacterial genomics have been mostly confined to single-species and single-task problems. For instance, genome-wide association studies (GWAS) have been successful in identifying genotype–phenotype associations within species, such as antimicrobial resistance or virulence traits (Lees et al., 2016; Power et al., 2017; San et al., 2020). More recent approaches such as unitig-based and deep learning models improved genotype–phenotype mapping by spanning the full pangenome (Lees et al., 2018; 2020) and predicting the effect of mutation based on the genomic context (Wiatrak et al., 2024). While the single-species models often perform well in their domains, they do not extend to other taxa and new genomic variants, and the huge genomic feature space relative to the number of labelled isolates leaves them prone to overfitting.

Genomic LMs. DNA LMs learn DNA sequence representations and have been shown to accurately represent long sequences (Zhou et al., 2023; Dalla-Torre et al., 2024; Jiang et al., 2023; Mourad, 2025; Nguyen et al., 2024; 2023), but are usually evaluated on human regulatory tasks, where transcriptional mechanisms and epigenomic regulation differ substantially from bacteria (Casadesús & Low, 2006). Moreover, even the DNA LMs with very large context window cannot span entire medium-sized bacterial genomes (Brix et al., 2025). **Protein LMs (pLMs)** learn representations that correlate with structure, stability, and function (Elnaggar et al., 2021; Lin et al., 2022; 2023; Hayes et al., 2025). As bacterial genomes consist largely of coding sequence and possess simpler regulatory architectures than eukaryotes, pLMs can capture a substantial fraction of relevant biology. pLMs model proteins in isolation, however, characterizing bacterial genomes requires modeling the contextual interactions between the proteins present in the genome. Finally, we differentiate a third group of genomic LMs - **Bacterial LMs (bLMs)** which are recently proposed genomic LMs that are purposefully built to model bacterial genomes. These include gLM2 (Cornman et al., 2024), a mixed-modality genomic LM that represents coding regions as amino acids and intergenic regions as nucleotides to model contiguous genome context, and Bacformer, a genome-level contextual protein LM that treats each bacterial genome as an ordered sequence of proteins, refining each protein vector in the presence of

¹<https://anonymous.4open.science/r/BacBench-B6EF>

Table 1: Summary of benchmarked models. “Max ctx.” = maximum context length supported at inference; “dim” = dimensionality of the output of the last hidden layer. DNA LMs, pLMs and bLMs are separated by a horizontal line.

Model	Input	Objective	Tokenisation	Params	dim	Training corpus	Max ctx.
Mistral-DNA (Mourad, 2025)	DNA	Autoregressive	Byte-pair	138M	768	Bacteria	512
DNABERT-2 (Zhou et al., 2023)	DNA	Masked	Byte-pair	117M	768	Multi-kingdom	512
Nucleotide Transformer (Dalla-Torre et al., 2024)	DNA	Masked	k -mer	250M	768	Multi-kingdom	2,048
ProkBERT (Ligeti et al., 2024)	DNA	Masked	k -mer	27M	384	Bacteria	4,096
Evo (Nguyen et al., 2024)	DNA	Autoregressive	Single nucleotide	6.5B	4,096	Multi-kingdom	8,192
ESM-2 (Lin et al., 2022)	Single protein seq.	Masked	Single amino acid	35M	480	Multi-kingdom	1,024
ESM-C (ESM Team, 2024)	Single protein seq.	Masked	Single amino acid	300M	960	Multi-kingdom	1,024
ProtBERT (Elnaggar et al., 2021)	Single protein seq.	Masked	Single amino acid	420M	1,024	Multi-kingdom	1,024
gLM2 (Corman et al., 2024)	Mixed modality (DNA & protein seq.)	Masked	Single nucleotide/amino acid	650M	1,280	Bacteria	4,096
Bacformer (Wiatrak et al., 2025)	Multiple protein seq.	Masked	Single protein	27M	480	Bacteria	6,000

all other proteins from the same genome, thus, encoding organism-level context. Both methods model the DNA or protein in the context of a bacterial genome and are pretrained on extensive bacterial corpora.

3 BACTERIAL GENOME REPRESENTATIONS & BASELINES

We selected a diverse set of five DNA LMs, three pLMs and two bLMs to represent bacterial genomes and evaluate their performance across distinct tasks and scales. These genomic LMs take as input DNA, proteins, or both (gLM2) and can therefore generalize across the bacterial tree of life. Moreover, the suite spans modalities (DNA-only, single-protein, mixed DNA-protein, genome-level protein), objectives (masked vs. autoregressive), and training corpora (bacteria-specific vs. cross-kingdom), enabling a controlled comparison of how context length, modality, and pretraining data shape genome embeddings (Table 1).

Bacterial genomes representations. For the gene- and system-level tasks with **DNA LMs** we embed the coding sequence plus upstream promoter of the gene; we split sequences longer than the model limit L into overlapping windows of length L and average their embeddings across all windows. For genome-level tasks, we tile each genome into chunks with overlap, embed each chunk, and average the results to extract a genome embedding (Appendix B). For the **pLMs** we embed each protein present in the bacterial genome independently and average its residue embeddings. To generate genome-scale representations, we calculate the mean of all protein vectors (Appendix B). Finally, for **bLMs** we use contiguous, genome-aware inputs: for **gLM2**, we feed mixed-modality genomic segments that encode coding regions as amino acids and intergenic regions as nucleotides, preserving local genome context across adjacent genes; for **Bacformer**, we represent each genome as an ordered sequence of proteins by obtaining per-protein tokens from its base pLM (ESM-2 35M), and pass these through a transformer to learn contextualized, genome-level protein embeddings.

4 PREDICTION TASKS AND EVALUATION RESULTS

In BacBench, we consider six tasks across three scales: (i) gene , (ii) system , and (iii) genome . At gene-level, we consider the task of gene essentiality. At system-level, we assess operon identification and protein-protein interaction prediction tasks. Finally, at genome-level we evaluate methods on strain clustering, antibiotic resistance and phenotypic traits prediction tasks. We briefly describe each task and include further experimental details in the Appendices A & B.

4.1 GENE ESSENTIALITY PREDICTION

Identifying essential genes is crucial for (i) defining the minimal set of cellular functions, and (ii) prioritizing drug targets. To distinguish essential from non-essential genes, the methods need to generalize to phylogenetically diverse bacteria beyond the species in the training set. For each model we report the performance by (i) fitting a linear classifier on top of the frozen gene embeddings, (ii) finetuning the model to predict a binary label (i.e. essentiality of input gene; Appendix B).

Data. We compile the dataset from the Database of Essential Genes (DEG) (Zhang et al., 2004), a hand-curated resource which aggregates studies published for a broad range of bacteria. After quality-control filtering, the corpus comprises 51 distinct genomes spanning 37 species (Appendix

A). This amounts to 22,486 essential and 146,922 non-essential genes. To prevent train-test leakage, we split by genus—placing all genomes from a genus in one split—and evaluate on held-out genera, enforcing generalization to phylogenetically distant strains.

Metrics. Gene essentiality prediction is a binary classification task. We evaluate performance using AUROC and AUPRC and report macro-average metrics across test genomes.

Results. pLMs and bLMs substantially outperform DNA-based models in terms of both AUROC (Fig. 2) and AUPRC using both linear probing and finetuning (Appendix C), suggesting that essential genes share conserved protein motifs that the protein-based embeddings capture. bLMs perform the best, with gLM2 achieving the best results overall, and Bacformer significantly outperforming its backbone ESM-2 model, showing the benefits of incorporating genome context. When considering DNA LMs only, Evo achieves the best performance, demonstrating the benefits of scaling model and context size.

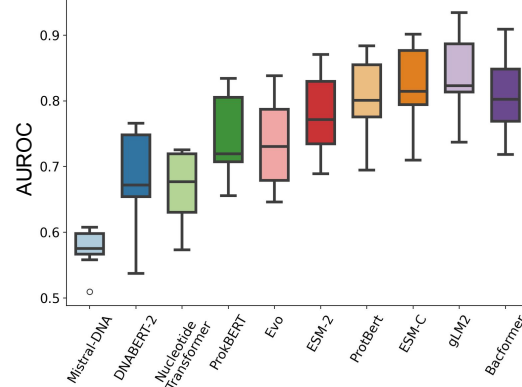


Figure 2: AUROC across genomes on essential gene prediction task using a linear model. The box spans the interquartile range with a line marking the median value.

4.2 OPERON IDENTIFICATION

Operons are multi-gene transcriptional units that underpin coordinated gene expression. Accurate operon maps are pivotal for (i) constructing gene regulatory networks (Fortino et al., 2014) and (ii) refining genome-scale metabolic models for strain engineering (Orth et al., 2010). In this task, the methods must predict whether the two neighbouring genes belong to the same operon, effectively predicting operon boundaries. Because available annotations are scarce we evaluate the task in a zero-shot setting.

Data. Due to the lack of experimentally validated operon annotations, we generated the operon labels by performing long-read RNA sequencing on a set of 5 diverse strains, amounting to 3,310 unique operons (Appendix A).

Metrics. In BacBench, operon identification is a zero-shot binary classification task. We use the cosine similarity value between the two genes combined with the information on the genes’ strand as a score indicating whether the genes belong to the same operon (Appendix A). We leverage AUROC and AUPRC for measuring performance. Finally, we report results across distinct strains.

Results. All methods except Bacformer attain similar performance (Fig. 3) in terms of both AUROC and AUPRC (Appendix C). We attribute the high performance of the Bacformer due to the (i) whole-genomic context of the model and (ii) extensive bacterial training corpus. ProkBERT performs 2nd best, showing that (i) scaling the model does not necessarily lead to improved results, (ii) the importance of a relevant pretraining corpus (ProkBERT was pretrained only on prokaryotes). Notably, both DNA and pLMs perform similarly on the task, indicating that the choice of modality is less important than incorporating genome-level context and domain-matched pretraining. Finally, the low overall performance across models highlights the need for task-specific methods and generating datasets which would allow for finetuning.

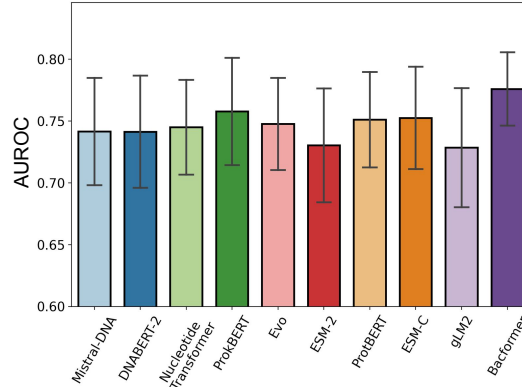


Figure 3: Zero-shot AUROC on operon identification task, the error bars represent standard error across strains.

4.3 PROTEIN-PROTEIN INTERACTION PREDICTION

Mapping protein-protein interactions (PPI) is central to (i) reconstructing gene networks (Snider et al., 2015) and (ii) prioritizing drug-target combinations (Wilson et al., 2022). Compared to the operon identification task where interacting genes lie next to each other, in the PPI task proteins can be separated by millions of base pairs. In this task each input is a pair of proteins, and the goal is to predict whether two proteins interact. Because interaction data are available at the protein-level, we restrict the benchmark to models which only require protein sequence data, specifically, pLMs and Bacformer. We evaluate two regimes: (i) zero-shot: the cosine similarity between the frozen embeddings of the two proteins serves as the interaction score, (ii) finetuned: the averaged embeddings of the two proteins are passed through a linear classifier trained to output a binary class.

Data. We downloaded and processed all 10,533 bacterial strains available in STRING DB (Szklarczyk et al., 2023) together with associated PPI scores. For every strain we extract the combined interaction score for each protein pair; the median strain contains almost 640,000 scored pairs. We binarize the labels, resulting in roughly 10% of all interactions being positive (Appendix A).

Metrics. PPI is a binary classification task. We report performance with AUROC and AUPRC, macro-averaged over test genomes. Both the positive and negative labels are provided by STRING.

Results. In both zero-shot and finetuned setups, Bacformer achieves substantially higher scores than other methods. We attribute this to the rotary positional embeddings (Su et al., 2024), which increase the cosine similarity score between neighbouring genes. This is in line with experimental data, which shows that the neighbouring genes in the bacterial genomes tend to interact with each other (Dandekar et al., 1998). We also notice how the difference in performance between the methods decreases following finetuning, demonstrating how the models can adapt to the task during training. Finally, the performance of ESM-2 compared to ESM-C and ProtBERT shows that scaling up the training data and model size does not benefit PPI prediction in bacteria.

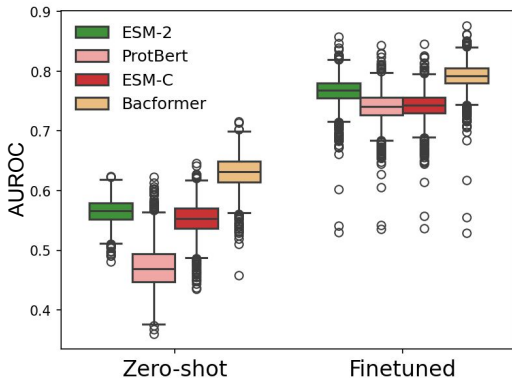


Figure 4: AUROC across genomes on PPI task in the zero-shot and finetuned setting. The box spans the interquartile range with a line marking the median value.

4.4 STRAIN CLUSTERING

Rapid clustering of whole genomes is valuable for placing newly sequenced metagenome assembled genomes (MAGs) into the bacterial tree of life and quality control. Therefore, we propose a metagenomic strain clustering task where the goal is to recover taxonomy using genome only. In this task, we feed every MAG to the model without using any species tokens or other metadata—so evaluation is fully zero-shot. We then evaluate whether models’ genome embeddings preserve the taxonomy. A good embedding should cluster genomes from the same species close together and, at broader levels, conserve members of e.g. the same genus or family. We perform the clustering by computing k nearest neighbors across different k and running Leiden clustering (Traag et al., 2019) at various resolutions (Appendix B).

Data. In this BacBench task, we draw 6,071 strains from MGnify (Mitchell et al., 2023), spanning 25 species distributed across 10 genera and 7 families chosen to give a balanced phylogenetic spread. We process DNA and protein inputs in the same way for every model; and average the vector from the final hidden layer over the whole genome to yield one fixed-length embedding per genome.

Metrics. We quantify clustering performance with the adjusted Rand index (ARI), normalized mutual information (NMI) and average silhouette width (ASW). Higher is better for all metrics. To obtain cluster assignments for each model we run Leiden clustering across resolutions, retaining the resolution that maximized the mean of the three metrics over the species, genus and family ranks. The ASW is unaffected by taxonomic rank, so we report it once in the *Combined* section.

Table 2: Clustering performance (higher is better) across taxonomic ranks and metrics. Best value for each metric is bolded. DNA LMs, pLMs and bLMs are separated by a horizontal line.

Method	Species		Genus		Family		Combined		
	ARI	NMI	ARI	NMI	ARI	NMI	ARI	NMI	ASW
Mistral-DNA	89.57	94.85	69.04	86.60	48.52	78.56	69.04	86.67	39.88
DNABERT-2	97.10	98.33	64.98	85.81	43.85	76.35	68.64	86.83	63.10
Nucleotide Transformer	98.14	98.89	64.06	85.28	43.14	75.84	68.45	86.67	39.24
ProkBERT	98.75	99.38	63.55	84.88	42.76	75.46	68.35	86.58	65.03
Evo	55.56	76.98	49.28	72.42	35.53	66.58	46.79	72.00	25.33
ESM-2	50.94	71.67	56.82	72.20	46.84	69.06	51.53	70.97	16.75
ESM-C	72.65	87.70	79.76	89.87	60.39	83.43	70.93	87.00	29.55
ProtBERT	68.53	86.47	85.98	92.40	65.78	86.42	73.43	88.43	39.12
gLM2	72.33	84.03	66.60	82.27	49.03	75.93	62.65	80.74	40.16
Bacformer	79.39	90.62	77.92	90.20	54.74	80.71	70.68	87.17	30.12

Results. At the species level, DNA LMs attain the highest scores, indicating that species boundaries can be recovered from sequence alone and consistent with the extra signal carried in non-coding regions. Moving up to genus and family levels, the advantage shifts: pLMs and Bacformer overtake the DNA LMs, suggesting that protein embeddings retain deeper evolutionary relationships more faithfully. Aggregated across all ranks, ProtBERT shows the highest ARI and NMI values, whereas the DNA models have the overall highest ASW. Interestingly, gLM2 which is a mixed modality model combining DNA and protein sequences performs worse than other models. Upon investigation, we believe this is due to the large variance in its embeddings and suggest that a mixed modality method with improved regularization could capture both fine-grained strain identity and higher-order phylogeny in a single embedding space.

4.5 ANTIBIOTIC RESISTANCE PREDICTION

Predicting antibiotic resistance is a task with (i) immediate clinical value for guiding antimicrobial drug therapy, (ii) monitoring the spread of resistant lineages and (iii) prioritizing compounds in drug-discovery pipelines. In this BacBench task, we define two subtasks: (1) given a bacterial genome, predict whether the strain is resistant or susceptible to a specific drug, and (2) given a bacterial genome, estimate its minimum inhibitory concentration (MIC). The first subtask is a binary classification problem (resistant vs susceptible), while the second subtask is a regression problem. Due to the computational complexity of computing genome-level embeddings on over 25k genomes (Appendix B), we (i) do not include Evo in the analysis due to its size (6.5B parameters), which makes it computationally infeasible to embed the entire corpus in most academic environments (see *Runtime analysis*; Appendix B), (ii) perform evaluation by stacking a linear layer on top of the frozen genome representation, and fine-tune a separate linear classifier for each model.

Data. We assemble a cross-species panel from the NIH Antimicrobial Susceptibility Test browser (National Center for Biotechnology Information, 2025), covering 25,032 strains drawn from 38 bacterial species. After quality control and removal of sparsely sampled drugs, the dataset retains 36 antibiotics for binary label and 56 for regression prediction that span diverse classes of antibiotics (Appendix A). For the binary task we use the resistant/susceptible calls and discard any ambiguous entries (Appendix A). For the regression task we extract the raw MIC values, apply a *log1p* transformation to dampen heavy tails and train a separate linear model for each drug-model pair.

Metrics. We report performance in the classification setting with AUROC and AUPRC across drugs. For the MIC regression we compute the *Pearson* correlation coefficient and the coefficient of determination (R^2) averaged across antibiotics. We report mean scores across antibiotics and include full per-antibiotic tables in the Appendix C.

Results. On both binary and regression setup, bLMs and pLMs tend to outperform DNA LMs. This may be explained by the fact that resistance is usually acquired through the mutation in the coding sequence, with studies showing that 90% of characterized resistance-conferring variants reside in coding regions (Sandgren et al., 2009; Farhat et al., 2019). Finally, the bLM-Bacformer achieves the best results implying the importance of epistatic effects on antibiotic resistance (Trindade et al., 2009) which the model considers by modeling the interactions between all of the proteins present in the genome. Notably, the methods record highly variable performance across antibiotics, underscoring the difficulty of building a single model that generalizes across the resistance mechanisms.

Table 3: Performance (% , higher is better) on the antibiotic-resistance prediction tasks across drugs. Values are mean $\pm$ standard deviation across 3 runs; best for each metric is bolded. DNA LMs, pLMs and bLMs are separated by a horizontal line.

Method	Binary		Regression	
	AUPRC	AUROC	R ²	Pearson
Mistral-DNA	49.86 $\pm$ 22.73	76.05 $\pm$ 9.14	19.71 $\pm$ 17.37	40.95 $\pm$ 21.51
DNABERT-2	52.35 $\pm$ 23.55	79.88 $\pm$ 7.26	23.61 $\pm$ 18.33	45.89 $\pm$ 19.27
Nucleotide Transformer	59.70 $\pm$ 22.90	84.22 $\pm$ 6.47	27.74 $\pm$ 17.80	51.19 $\pm$ 16.76
ProkBERT	58.90 $\pm$ 23.67	83.94 $\pm$ 6.39	28.00 $\pm$ 17.57	51.24 $\pm$ 16.91
ESM-2	62.04 $\pm$ 23.64	85.05 $\pm$ 6.63	31.18 $\pm$ 17.39	54.60 $\pm$ 16.15
ESM-C	63.41 $\pm$ 23.43	85.68 $\pm$ 6.81	33.15 $\pm$ 17.62	56.79 $\pm$ 15.46
ProtBERT	61.79 $\pm$ 23.19	84.51 $\pm$ 6.31	28.23 $\pm$ 18.46	51.56 $\pm$ 17.61
gLM2	55.80 $\pm$ 24.46	82.16 $\pm$ 6.75	26.15 $\pm$ 18.13	49.41 $\pm$ 17.14
Bacformer	67.97 $\pm$ 20.73	87.61 $\pm$ 5.68	33.84 $\pm$ 18.60	57.53 $\pm$ 15.43

4.6 PHENOTYPIC TRAITS PREDICTION

Accurately predicting phenotype from a genomic sequence enables (i) inference of the biological or ecological function of bacteria (Feldbauer et al., 2015) and (ii) engineering organisms with the exact metabolic traits needed for efficient waste remediation (Rafeeq et al., 2023), accelerating sustainable industrial bioproduction (Lawson et al., 2021). We evaluate whether the models’ genome embeddings are predictive of diverse phenotypic traits. Here, given a genome embedding, the task is to predict a trait. Similarly as in the antibiotic resistance prediction task above, we perform linear probing evaluation, training a separate linear classifier per phenotype and exclude the Evo model due to the computational costs (Appendix B).

Data. To create the benchmark we collated large trait inventories (Madin et al., 2020; Brbić et al., 2016; Weimann et al., 2016), apply stringent quality filters and discard traits represented by only a handful of isolates (Appendix A). The final corpus covers 139 discrete phenotypes spanning 15, 477 bacterial species, making it the broadest dataset of its kind and challenging the models to generalize well beyond their training clades. We group traits into carbon utilisation, biochemical activity, growth conditions and cellular morphology (Appendix A), which we use later for stratified analysis. As similar genomes often share the same phenotype, we split the data for each phenotype by genus—placing all genomes from a genus in one split—and evaluate on held-out genera, enforcing generalization to phylogenetically distant strains.

Metrics. For BacBench, we restrict the phenotypic traits to categorical traits due to their sufficient number of available samples, and evaluate performance using the macro-averaged AUROC and AUPRC over phenotype categories. We report mean scores for every phenotype group and include full per-trait tables in the Appendix C.

Results. pLMs and bLMs outperform DNA LMs across phenotype groups and metrics (Fig. 5 & Appendix C). The relative difference is especially large for biochemical activity and carbon utilization traits, likely because these phenotypes hinge on enzyme active-site composition and pathway membership—information that is explicit in amino-acid space and can be better captured

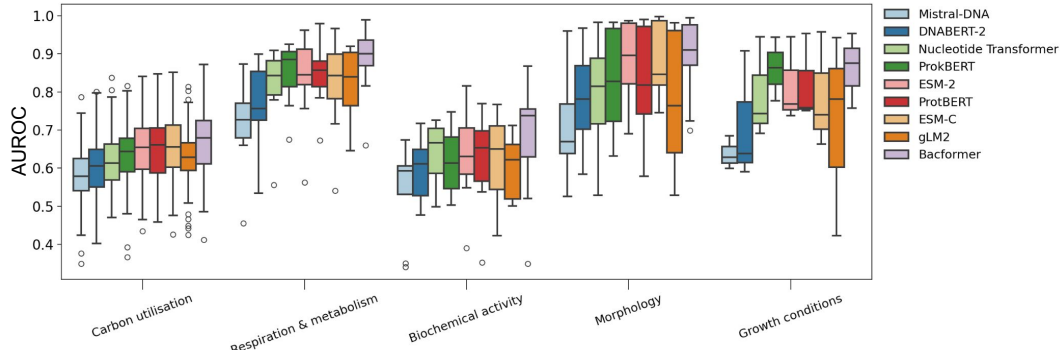


Figure 5: AUROC across diverse phenotypic traits groups and methods. The box spans the inter-quartile range with a line marking the median value.

by pLMs (Teukam et al., 2024; Lawson et al., 2021). The Bacformer bLM attains the highest scores, indicating that whole-genome context helps explain phenotypes arising from coordinated protein function and epistatic interactions (Trindade et al., 2009). In contrast, Mistral-DNA, the only autoregressive model in the task, lags well behind the other models. Overall performance remains moderate—especially for carbon utilization, biochemical activity and growth condition traits—highlighting considerable room for improvement. Progress will likely require more complex contextualized models that incorporate environmental metadata, together with larger, more balanced phenotype datasets.

5 DISCUSSION

Summary. BacBench addresses the lack of a comprehensive, cross-species evaluation resource for bacterial genomics by providing unified datasets and benchmarks, and by evaluating existing genomic language models across species, tasks, and biological scales. It introduces a newly generated dataset for operon identification—a key problem for refining genome-scale metabolic models—and curates five additional tasks (gene essentiality, protein–protein interaction, strain clustering, antibiotic resistance, and 139 phenotypic traits) into a single framework covering 67k genomes from 17.6k species. Our experiments show that (i) existing genomic LMs (DNA LMs, pLMs, and bLMs) capture core taxonomic structure and functional relatedness, providing a strong baseline representation of bacterial genomes, but fall short at accurate genome-to-phenotype prediction, as evidenced by antibiotic resistance and phenotype tasks; (ii) models purpose-built for bacteria (gLM2, Bacformer) or trained on bacteria-specific corpora (ProkBERT) tend to outperform broad, cross-kingdom counterparts across most tasks, underscoring the value of domain-matched pretraining and inductive biases; (iii) different modeling approaches excel at different problems—DNA LMs capture fine-grained taxonomic signals and do well on operon identification, pLMs better preserve deeper phylogeny and functional similarity for strain clustering, and bLMs like Bacformer perform best on tasks driven by multi-gene interactions (phenotypes, antibiotic resistance). All datasets are accompanied with extensive documentation; the embedding and evaluation library is provided at <https://anonymous.4open.science/r/BacBench-B6EF>.

Towards accurate genome-to-phenotype bacterial prediction. Our results suggest a practical path forward towards building model for genome-to-phenotype mapping in bacteria: (i) the model should represent entire genomes end-to-end to capture long-range, cross-protein dependencies (currently only feasible in Bacformer); (ii) pretrain on substantially larger, bacteria-focused corpora (we estimate >4M unique strains remain untapped (Mitchell et al., 2023; Markowitz et al., 2012; Blackwell et al., 2024)); (iii) integrate DNA and protein modalities—and where available RNA—to couple regulatory and coding signals in a single embedding while remaining computationally efficient; (iv) allow inclusion of structured priors (e.g., resistance gene catalogs, operon maps, HGT markers) to improve data efficiency on scarce phenotype labels; (v) expand high-quality phenotype supervision (including knock-outs and standardized trait panels) to close the supervision gap limiting genome-to-phenotype learning.

Limitations & Future Work. By releasing BacBench we provide a foundation for more expressive, cross-species models of bacterial genomics, yet the present benchmark covers only part of the functional landscape. Other tasks are not yet included, and certain phenotypes and antibiotic classes remain sparse (Appendix A), underscoring the need for generating new data. We benchmarked a representative set of publicly available models capable of cross-species generalization, and expect the model suite to expand in future iterations. Finally, the current tasks do not include modalities such as transcriptomics and metabolomics due to data sparsity and inconsistent metadata. As community datasets mature, incorporating multi-omics will enable more faithful evaluation of causal, context-dependent genome-to-phenotype mappings.

We anticipate that subsequent iterations will broaden task and model coverage, ultimately enabling contextual, genome-scale representations for bacterial genomes. Community contributions of datasets, models, and evaluation routines are encouraged so that BacBench evolves into a continually updated standard for bacterial ML.

6 REPRODUCIBILITY STATEMENT

We release an *anonymous* codebase for preprocessing, embedding, and evaluation at <https://anonymous.4open.science/r/BacBench-B6EF>, together with helper utilities for bacterial genomics to lower the barrier for adding new models and tasks. All datasets in BacBench are fully documented and accompanied by anonymized MLCommons Croissant metadata files in the supplementary materials, enabling unambiguous data loading and lineage tracking. The Appendix details quality filtering, preprocessing pipelines, train/validation/test splits, model and training configurations, and hyperparameters; it also specifies random seeds, hardware, and optimization settings. Further details on the exact scripts to reproduce the results can be found in the anonymous repository linked above. These artifacts together provide the necessary pointers—code, data descriptors, and experimental specifications—to reproduce our results and to extend BacBench as an evolving benchmark for the ML community.

REFERENCES

- Shohreh Ariaeenejad, Javad Gharechahi, Mehdi Foroozandeh Shahraki, Fereshteh Fallah Atanaki, Jian-Lin Han, Xue-Zhi Ding, Falk Hildebrand, Mohammad Bahram, Kaveh Kavousi, and Ghasem Hosseini Salekdeh. Precision enzyme discovery through targeted mining of metagenomic data. *Natural Products and Bioprospecting*, 14(1):7, 2024.
- Jimmy Lei Ba, Jamie Ryan Kiros, and Geoffrey E Hinton. Layer normalization. *arXiv preprint arXiv:1607.06450*, 2016.
- Grace Blackwell, Karel M. Brinda, John A. Lees, Rebecca A. Gladstone, Philip Brownridge, et al. Allthebacteria: a uniformly assembled and searchable compendium of all public bacterial genomes. *bioRxiv*, pp. 2024.03.08.584059, 2024. doi: 10.1101/2024.03.08.584059.
- Maria Brbić, Matija Piškorec, Vedrana Vidulin, Anita Kriško, Tomislav Šmuc, and Fran Supek. The landscape of microbial phenotypic traits and associated genes. *Nucleic acids research*, pp. gkw964, 2016.
- Garyk Brixi, Matthew G Durrant, Jerome Ku, Michael Poli, Greg Brockman, Daniel Chang, Gabriel A Gonzalez, Samuel H King, David B Li, Aditi T Merchant, et al. Genome modeling and design across all domains of life with evo 2. *BioRxiv*, pp. 2025–02, 2025.
- Josep Casadesús and David Low. Epigenetic gene regulation in bacteria. *Molecular Microbiology*, 60(4):820–826, 2006. doi: 10.1111/j.1365-2958.2006.05136.x.
- Gareth A Coleman, Adrián A Davín, Tara A Mahendrarajah, Lénárd L Szánthó, Anja Spang, Philip Hugenholtz, Gergely J Szöllősi, and Tom A Williams. A rooted phylogeny resolves early bacterial evolution. *Science*, 372(6542):eabe0511, 2021.
- CRyPTIC Consortium. A data compendium associating the genomes of 12,289 mycobacterium tuberculosis isolates with quantitative resistance phenotypes to 13 antibiotics. *PLoS biology*, 20(8): e3001721, 2022.
- Andre Cornman, Jacob West-Roberts, Antonio Pedro Camargo, Simon Roux, Martin Beracochea, Milot Mirdita, Sergey Ovchinnikov, and Yunha Hwang. The omg dataset: An open metagenomic corpus for mixed-modality genomic language modeling. *bioRxiv*, pp. 2024–08, 2024.
- Hugo Dalla-Torre, Liam Gonzalez, Javier Mendoza-Revilla, Nicolas Lopez Carranza, Adam Henryk Grzywaczewski, Francesco Oteri, Christian Dallago, Evan Trop, Bernardo P de Almeida, Hassan Sirelkhatim, et al. Nucleotide transformer: building and evaluating robust foundation models for human genomics. *Nature Methods*, pp. 1–11, 2024.
- Thomas Dandekar, Berend Snel, Martijn Huynen, and Peer Bork. Conservation of gene order: a fingerprint of proteins that physically interact. *Trends in biochemical sciences*, 23(9):324–328, 1998.

- Petr Danecek, James K Bonfield, Jennifer Liddle, John Marshall, Valeriu Ohan, Martin O Pollard, Andrew Whitwham, Thomas Keane, Shane A McCarthy, Robert M Davies, and Heng Li. Twelve years of SAMtools and BCFtools. *GigaScience*, 10(2), 02 2021. ISSN 2047-217X. doi: 10.1093/gigascience/giab008. URL <https://doi.org/10.1093/gigascience/giab008>. giab008.
- Wouter AA de Steenhuijsen Piters, Elisabeth AM Sanders, and Debby Bogaert. The role of the local microbial ecosystem in respiratory health and disease. *Philosophical Transactions of the Royal Society B: Biological Sciences*, 370(1675):20140294, 2015.
- Ahmed Elnaggar, Michael Heinzinger, Christian Dallago, Ghalia Rehawi, Yu Wang, Llion Jones, Tom Gibbs, Tamas Feher, Christoph Angerer, Martin Steinegger, et al. Prottrans: Toward understanding the language of life through self-supervised learning. *IEEE transactions on pattern analysis and machine intelligence*, 44(10):7112–7127, 2021.
- ESM Team. Esm cambrian: Revealing the mysteries of proteins with unsupervised learning. <https://evolutionaryscale.ai/blog/esm-cambrian>, 2024. EvolutionaryScale website, published 4 December 2024, accessed 17 April 2025.
- Maha R Farhat, Luca Freschi, Roger Calderon, Thomas Ioerger, Matthew Snyder, Conor J Meehan, Bouke de Jong, Leen Rigouts, Alex Sloutsky, Devinder Kaur, et al. Gwas for quantitative resistance phenotypes in mycobacterium tuberculosis reveals resistance genes and regulatory regions. *Nature communications*, 10(1):2128, 2019.
- Roman Feldbauer, Frederik Schulz, Matthias Horn, and Thomas Rattei. Prediction of microbial phenotypes based on comparative genomics. *BMC bioinformatics*, 16:1–8, 2015.
- Vittorio Fortino, Olli-Pekka Smolander, Petri Auvinen, Roberto Tagliaferri, and Dario Greco. Transcriptome dynamics-based operon prediction in prokaryotes. *BMC bioinformatics*, 15:1–14, 2014.
- Thomas Hayes, Roshan Rao, Halil Akin, Nicholas J Sofroniew, Deniz Oktay, Zeming Lin, Robert Verkuil, Vincent Q Tran, Jonathan Deaton, Marius Wiggert, et al. Simulating 500 million years of evolution with a language model. *Science*, pp. eads0018, 2025.
- Albert Q. Jiang, Alexandre Sablayrolles, Arthur Mensch, Chris Bamford, Devendra Singh Chaplot, Diego de las Casas, Florian Bressand, Gianna Lengyel, Guillaume Lample, Lucile Saulnier, L  lio Renard Lavaud, Marie-Anne Lachaux, Pierre Stock, Teven Le Scao, Thibaut Lavril, Thomas Wang, Timoth  e Lacroix, and William El Sayed. Mistral 7b, 2023. URL <https://arxiv.org/abs/2310.06825>.
- Diederik P Kingma. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*, 2014.
- Jenna Morgan Lang, Aaron E Darling, and Jonathan A Eisen. Phylogeny of bacterial and archaeal genomes using conserved genes: supertrees and supermatrices. *PloS one*, 8(4):e62510, 2013.
- Christopher E Lawson, Jose Manuel Mart  , Tijana Radivojevic, Sai Vamshi R Jonnalagadda, Reinhard Gentz, Nathan J Hillson, Sean Peisert, Joonhoon Kim, Blake A Simmons, Christopher J Petzold, et al. Machine learning for metabolic engineering: A review. *Metabolic Engineering*, 63:34–60, 2021.
- John A. Lees, Minna Vehkala, Niko V  lim  ki, Simon R. Harris, and Claire *et al.* Chewapreecha. Sequence element enrichment analysis to determine the genetic basis of bacterial phenotypes. *Nature Communications*, 7:12797, 2016. doi: 10.1038/ncomms12797.
- John A Lees, Marco Galardini, Stephen D Bentley, Jeffrey N Weiser, and Jukka Corander. Pyseer: a comprehensive tool for microbial pangenome-wide association studies. *Bioinformatics*, 34(24):4310–4312, 2018.
- John A Lees, T Tien Mai, Marco Galardini, Nicole E Wheeler, Samuel T Horsfield, Julian Parkhill, and Jukka Corander. Improved prediction of bacterial genotype-phenotype associations using interpretable pangenome-spanning regressions. *MBio*, 11(4):10–1128, 2020.

- Heng Li. Minimap2: pairwise alignment for nucleotide sequences. *Bioinformatics*, 34(18):3094–3100, 2018.
- Balázs Ligeti, István Szepesi-Nagy, Babett Bodnár, Noémi Ligeti-Nagy, and János Juhász. Prokbert family: genomic language models for microbiome applications. *Frontiers in Microbiology*, 14, 2024. ISSN 1664-302X. doi: 10.3389/fmicb.2023.1331233. URL <https://www.frontiersin.org/articles/10.3389/fmicb.2023.1331233>.
- Zeming Lin, Halil Akin, Roshan Rao, Brian Hie, Zhongkai Zhu, Wenting Lu, Nikita Smetanin, Allan dos Santos Costa, Maryam Fazel-Zarandi, Tom Sercu, Sal Candido, et al. Language models of protein sequences at the scale of evolution enable accurate structure prediction. *bioRxiv*, 2022.
- Zeming Lin, Halil Akin, Roshan Rao, Brian Hie, Zhongkai Zhu, Wenting Lu, Nikita Smetanin, Robert Verkuil, Ori Kabeli, Yaniv Shmueli, et al. Evolutionary-scale prediction of atomic-level protein structure with a language model. *Science*, 379(6637):1123–1130, 2023.
- Zhaowei Luo, Zhanghua Qi, Jie Luo, and Tingtao Chen. Potential applications of engineered bacteria in disease diagnosis and treatment. *Microbiome Research Reports*, 4(1):N–A, 2024.
- Joshua S Madin, Daniel A Nielsen, Maria Brbic, Ross Corkrey, David Danko, Kyle Edwards, Martin KM Engqvist, Noah Fierer, Jemma L Geoghegan, Michael Gillings, et al. A synthesis of bacterial and archaeal phenotypic trait data. *Scientific data*, 7(1):170, 2020.
- Victor M Markowitz, I-Min A Chen, Krishna Palaniappan, Ken Chu, Ernest Szeto, Yuri Grechkin, Anna Ratner, Biju Jacob, Jinghua Huang, Peter Williams, et al. IMG: the integrated microbial genomes database and comparative analysis system. *Nucleic acids research*, 40(D1):D115–D122, 2012.
- Alex L. Mitchell, Antonio Almeida, Martin Beracochea, Matthew Boland, Jack Burgin, Guy Cochrane, et al. Mgnify: the microbiome analysis resource in 2023. *Nucleic Acids Research*, 51(D1):D723–D730, 2023. doi: 10.1093/nar/gkac1088.
- Raphael Mourad. Mistral-DNA. <https://github.com/raphaelmourad/Mistral-DNA>, 2025. GitHub repository, accessed 17 April 2025.
- National Center for Biotechnology Information. Antibiotic susceptibility test (ast) browser, 2025. URL <https://www.ncbi.nlm.nih.gov/pathogens/ast>. Accessed 24 Apr 2025.
- Eric Nguyen, Michael Poli, Marjan Faizi, Armin Thomas, Michael Wornow, Callum Birch-Sykes, Stefano Massaroli, Aman Patel, Clayton Rabideau, Yoshua Bengio, et al. Hyenadna: Long-range genomic sequence modeling at single nucleotide resolution. *Advances in neural information processing systems*, 36:43177–43201, 2023.
- Eric Nguyen, Michael Poli, Matthew G Durrant, Brian Kang, Dhruva Katrekar, David B Li, Liam J Bartie, Armin W Thomas, Samuel H King, Garyk Brixi, et al. Sequence modeling and design from molecular to genome scale with evo. *Science*, 386(6723):ead09336, 2024.
- Jeffrey D Orth, Ines Thiele, and Bernhard Ø Palsson. What is flux balance analysis? *Nature biotechnology*, 28(3):245–248, 2010.
- A Paszke. Pytorch: An imperative style, high-performance deep learning library. *arXiv preprint arXiv:1912.01703*, 2019.
- Robert A. Power, Julian Parkhill, and Tulio de Oliveira. Microbial genome-wide association studies: lessons from human gwas. *Nature Reviews Genetics*, 18:41–50, 2017. doi: 10.1038/nrg.2016.132.
- Hamza Rafeeq, Nadia Afsheen, Sadia Rafique, Arooj Arshad, Maham Intisar, Asim Hussain, Muhammad Bilal, and Hafiz MN Iqbal. Genetically engineered microorganisms for environmental remediation. *Chemosphere*, 310:136751, 2023.
- James Emmanuel San, Shakuntala Baichoo, Aquillah Kanzi, Yumna Moosa, Richard Lessells, Vagner Fonseca, John Mogaka, Robert Power, and Tulio de Oliveira. Current affairs of microbial genome-wide association studies: approaches, bottlenecks and analytical pitfalls. *Frontiers in microbiology*, 10:3119, 2020.

- Andreas Sandgren, Michael Strong, Preetika Muthukrishnan, Brian K Weiner, George M Church, and Megan B Murray. Tuberculosis drug resistance mutation database. *PLoS medicine*, 6(2):e1000002, 2009.
- Célio Dias Santos-Júnior, Marcelo DT Torres, Yiqian Duan, Álvaro Rodríguez Del Río, Thomas SB Schmidt, Hui Chong, Anthony Fullam, Michael Kuhn, Chengkai Zhu, Amy Houseman, et al. Discovery of antimicrobial peptides in the global microbiome with machine learning. *Cell*, 187(14):3761–3778, 2024.
- Jamie Snider, Max Kotlyar, Punit Saraon, Zhong Yao, Igor Jurisica, and Igor Stagljar. Fundamentals of protein interaction network mapping. *Molecular systems biology*, 11(12):848, 2015.
- Jianlin Su, Murtadha Ahmed, Yu Lu, Shengfeng Pan, Wen Bo, and Yunfeng Liu. Roformer: Enhanced transformer with rotary position embedding. *Neurocomputing*, 568:127063, 2024.
- Damian Szklarczyk, Rebecca Kirsch, Mikaela Koutrouli, Katerina Nastou, Farrokh Mehryary, Radja Hachilif, Annika L Gable, Tao Fang, Nadezhda T Doncheva, Sampo Pyysalo, et al. The string database in 2023: protein–protein association networks and functional enrichment analyses for any sequenced genome of interest. *Nucleic acids research*, 51(D1):D638–D646, 2023.
- Yves Gaetan Nana Teukam, Loïc Kwate Dassi, Matteo Manica, Daniel Probst, Philippe Schwaller, and Teodoro Laino. Language models can identify enzymatic binding sites in protein sequences. *Computational and Structural Biotechnology Journal*, 23:1929–1937, 2024.
- Vincent A Traag, Ludo Waltman, and Nees Jan Van Eck. From louvain to leiden: guaranteeing well-connected communities. *Scientific reports*, 9(1):1–12, 2019.
- Sandra Trindade, Ana Sousa, Karina Bivar Xavier, Francisco Dionisio, Miguel Godinho Ferreira, and Isabel Gordo. Positive epistasis drives the acquisition of multidrug resistance. *PLoS genetics*, 5(7):e1000578, 2009.
- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. *Advances in neural information processing systems*, 30, 2017.
- Saidi Wang, Minerva Ventolero, Haiyan Hu, and Xiaoman Li. A revisit to universal single-copy genes in bacterial genomes. *Scientific Reports*, 12(1):14550, 2022. doi: 10.1038/s41598-022-18762-z.
- Aaron Weimann, Kyra Mooren, Jeremy Frank, Phillip B Pope, Andreas Bremges, and Alice C McHardy. From genomes to phenotypes: Traitair, the microbial trait analyzer. *MSystems*, 1(6):10–1128, 2016.
- Jacob West-Roberts, Joshua Kravitz, Nishant Jha, Andre Cornman, and Yunha Hwang. Diverse genomic embedding benchmark for functional evaluation across the tree of life. *bioRxiv*, pp. 2024–07, 2024.
- Maciej Wiatrak, Aaron Weimann, Adam Dinan, Maria Brbić, and A. Floto, R. Sequence-based modelling of bacterial genomes enables accurate antibiotic resistance prediction. *bioRxiv*, pp. 2024.01.03.574022, 2024. doi: 10.1101/2024.01.03.574022.
- Maciej Wiatrak, Ramon Viñas Torné, Maria Ntemourtsidou, Adam Dinan, David C Abelson, Divya Arora, Maria Brbić, Aaron Weimann, and R Andres Floto. A contextualised protein language model reveals the functional syntax of bacterial evolution. *bioRxiv*, pp. 2025–07, 2025.
- Jennifer L Wilson, Ethan Steinberg, Rebecca Raczy, Russ B Altman, Nigam Shah, and Kevin Grimes. A network paradigm predicts drug synergistic effects using downstream protein–protein interactions. *CPT: Pharmacometrics & Systems Pharmacology*, 11(11):1527–1538, 2022.
- F Alexander Wolf, Philipp Angerer, and Fabian J Theis. Scanpy: large-scale single-cell gene expression data analysis. *Genome biology*, 19:1–5, 2018.
- Xihui Xu and Jiandong Jiang. Engineering microbiomes for enhanced bioremediation. *PLoS biology*, 22(12):e3002951, 2024.

Ren Zhang, Hong-Yu Ou, and Chun-Ting Zhang. Deg: a database of essential genes. *Nucleic acids research*, 32(suppl.1):D271–D272, 2004.

Zhihan Zhou, Yanrong Ji, Weijian Li, Pratik Dutta, Ramana Davuluri, and Han Liu. Dnabert-2: Efficient foundation model and benchmark for multi-species genome. *arXiv preprint arXiv:2306.15006*, 2023.

A ADDITIONAL DATASETS & TASKS DETAILS

In this section we outline the dataset details, including the overall and per dataset statistics as well as the preprocessing details. For each dataset we performed additional quality checks which we detail below. All of the genomes across datasets contain genome ID, which can be used to identify the genome source. Finally, we provide metadata on genes, proteins and genomes together with the datasets when available, together with extensive documentation.

Supplementary Table 1: Dataset summary for the six benchmark tasks used in this study.

Task	# Species	# Genomes	# Proteins	# Base pairs
Gene essentiality prediction	37	51	169 k	279 MB
Operon identification	5	5	22 k	25 MB
Protein–protein interaction prediction	6 956	10 533	36 M	N/A
Strain clustering	25	6 710	14 M	16 GB
Phenotypic traits prediction	15 477	24 462	100 M	111 GB
Antibiotic resistance prediction	38	26 302	105 M	112 GB

A.1 GENE ESSENTIALITY PREDICTION

We downloaded the gene essentiality annotations for bacteria across genomes from the Database of Essential Genes (DEG, <http://origin.tubic.org/deg>) (Zhang et al., 2004). Using the genome RefSeq ID provided in the database, we downloaded the associated genomes in both DNA and protein sequence modalities from the NCBI GenBank (<https://www.ncbi.nlm.nih.gov/genbank/>). Across 66 genomes from DEG, there were multiple genomes with more than 98% overlap when it comes to annotations. We therefore removed these genomes, as including it could lead to inflated evaluation metrics, leaving us with 51 genomes across 37 distinct species. For each genome we provide start and end for each gene together with essentiality annotations (*Yes*=essential, *No*=non-essential), verifying the gene locations are correct. We also provide the strand of the gene to allow for the extraction of the region upstream of the gene.

Split. We performed a random data split into training, validation, and test sets in a 60 / 20 / 20 % ratio. Additionally, to prevent train-test leakage, we split by genus—placing all genomes from a genus in one split—and evaluate on held-out genera, enforcing generalization to phylogenetically distant strains. AUROC

A.2 OPERON IDENTIFICATION

Due to the lack of reliable whole-genome operon annotations, we performed long-read RNA sequencing to annotate operons across five distinct strains, processing three independent biological replicates for each strain.

Bacterial strains and culture conditions. Five bacterial strains were used in this experiment: *Staphylococcus aureus* RN450 (*S. aureus* RN450), *Mycobacterium abscessus* ATCC 19977 (*M. ab*), Δ leuD Δ panCD *Mycobacterium tuberculosis* H37Rv 102J23 (Δ leuD Δ panCD *M. tb*), *Pseudomonas aeruginosa* PAO1, and *Escherichia coli* DH5 α . Each strain was cultured in triplicate under nutrient-rich conditions until mid-exponential phase ($OD_{600} = 0.4$ – 0.6) was reached.

RNA isolation. Cells were harvested by centrifugation and processed for total RNA extraction using the MasterPure Complete DNA and RNA Purification Kit (Lucigen) with strain-specific modifications: For *S. aureus* RN450, cell pellets were pre-treated with lysostaphin (Tris buffer, pH 8.0) at 37 °C

for 30 min to aid lysis. For the mycobacterial strains (*M. ab* and $\Delta\text{leuD } \Delta\text{panCD } M. tb$), cells were mechanically disrupted by bead beating in lysis buffer, followed by extraction with the standard kit protocol. Isolated RNA was treated twice with TURBO DNase (Invitrogen) to remove residual genomic DNA and purified using RNA Clean & Concentrator columns (Zymo Research). RNA integrity was assessed on an Agilent TapeStation RNA ScreenTape system, and concentrations were measured with a Qubit fluorometer (Invitrogen).

Library preparation and sequencing. For each replicate, 1,000 ng of total RNA underwent rRNA depletion with riboPOOLs (siTOOLs Biotech). The depleted RNA was polyadenylated with poly(A) polymerase (PAP) in the presence of a manganese catalyst, adding 50–90 adenosines per molecule. cDNA libraries were prepared with the Nanopore cDNA-PCR kit, pooled and sequenced on a PromethION device equipped with R10 flow cells (Oxford Nanopore Technologies).

Long-read RNA sequencing data preprocessing. For each sequenced strain, the following genome assemblies and gene annotations from NCBI RefSeq [ref] were used: GCF_000013425.1 (*Staphylococcus aureus* RN450), GCF_000069185.1 (*Mycobacterium abscessus* ATCC 19977), GCF_000195955.2 ($\Delta\text{leuD } \Delta\text{panCD } \textit{Mycobacterium tuberculosis} H37Rv 102J23), GCF_000006765.1 - (*Pseudomonas aeruginosa* PAO1), and GCF_002899475.1 (*Escherichia coli* DH5 α).$

ONT reads from each replicate were polished with Pypochopper (v2.7.10), polyA tails longer than 10 bases and sequencing adapters were trimmed using cutadapt and mapped against the genome assemblies using Minimap2 (v2.29)(Li, 2018) and Samtools (v1.22) (Danecek et al., 2021). Candidate operons were identified from read alignments spanning at least two genes on the same strand and then extended by combining overlapping candidates at most 50 base pairs apart. Operons were then collated from the triplicates for each strain and used as our operon annotations.

Split. We evaluate the operon identification in a zero-shot manner, therefore, we do not split the data into train, validation and test splits and use the entire dataset as a test set.

A.3 PROTEIN-PROTEIN INTERACTION PREDICTION

We downloaded all the data from the STRING DB download site (<https://string-db.org/cgi/download>). Using the species metadata file we selected only bacterial organisms and downloaded the protein sequences for them together with protein-protein interaction scores for protein pairs. After running the download scripts we ended up with 10,533 unique strains across 6,956 species. We used the *combined* interaction score which combines information from various sources to get a final score. STRING DB provides only protein sequences and no DNA, and the interaction scores are computed mainly at the protein-level, therefore, for this dataset we only provide protein sequences and omit DNA. To binarize the interaction scores, we set the threshold at 0.6 (≥ 0.6 =interaction, <0.6 =no interaction). This threshold was chosen through conducting small-scale experiments and looking at the average performance of AUROC and AUPRC on the validation set across genomes, choosing the threshold which attains the best average performance across the two.

Split. We performed a random data split into training, validation, and test sets in a 70 / 10 / 20 % ratio. The larger proportion of the train set compared to the gene essentiality task is motivated by the larger size of the overall dataset, with the 10% validation set still allowing for meaningful evaluation.

A.4 STRAIN CLUSTERING

To extract the metagenome assembled genomes (MAGs) for strain clustering, we use MGnify (Mitchell et al., 2023), which is a large-scale bacterial genomics database containing a diverse set of MAGs across numerous environments. The main reason for choosing MGnify over other potential resources is its large size combined with a uniform processing pipeline, providing comparable genomes. We wanted to evaluate whether various methods capture phylogenetic similarities across different taxonomic levels, therefore, we looked at strains which span different species, genera and families. These nested ranks provide three increasingly coarse resolutions to test whether an embedding preserves evolutionary signal. We use the taxonomic annotations provided by MGnify. The total number of unique genomes in the dataset is 6,071.

To extract the genomes of interest, we extracted the most common species from MGnify from the corpus of 300k bacterial genomes and selected 25 species which are distributed across distinct genera and families. For a meaningful evaluation each genus (or family) must contain *at least two* species, otherwise the genus- or family-level clustering metrics become degenerate (every strain would be trivially assigned to a unique cluster at that rank). We download the chosen assemblies and use the accompanying annotations to translate gene DNA to protein sequences, while retaining the original DNA for the DNA-modality experiments.

Split. The strain clustering task is a fully unsupervised task, therefore, we do not split the data into train, validation and test splits and use the entire dataset as a test set.

A.5 ANTIBIOTIC RESISTANCE PREDICTION

We leveraged the antibiotic sensitivity readings from the NCBI AST browser (<https://www.ncbi.nlm.nih.gov/pathogens/ast>), which contains hundreds of thousands of antibiotic-susceptibility test (AST) records for a diverse set of antibiotics. Using the genome assembly identifiers provided by the NCBI AST browser, we downloaded the DNA and protein sequences for each genome from the NCBI GenBank (<https://www.ncbi.nlm.nih.gov/genbank/>) and matched them with the antibiotic resistance readings. This left us with 26,052 unique genomes with matched antibiotic resistance labels. We then processed the antibiotic resistance readings into (i) binary and (ii) regression labels described below.

Binary. For binary prediction, antibiotic sensitivity labels from NCBI AST browser of either *sensitive* (S) or *resistant* (R) were used for training and testing. If the antibiotic sensitivity test had no S/R label, they were not included. We remove antibiotics which 1) have less than 500 available unique genomes in total, and 2) have less than 50 unique genomes per class. This is motivated by the need to ensure that every classifier is trained on a sufficiently large and reasonably balanced data set; with fewer than 500 genomes overall, or fewer than 50 genomes in either class, the resulting model would suffer from poor statistical power and unreliable performance estimates. This resulted in 37 unique antibiotics. The Supp. Table 2 shows the number of available genomes per drug, including the number of susceptible and resistant genomes. The number of available readings varies strongly between antibiotics, which partly explains the high variance between the per drug performance.

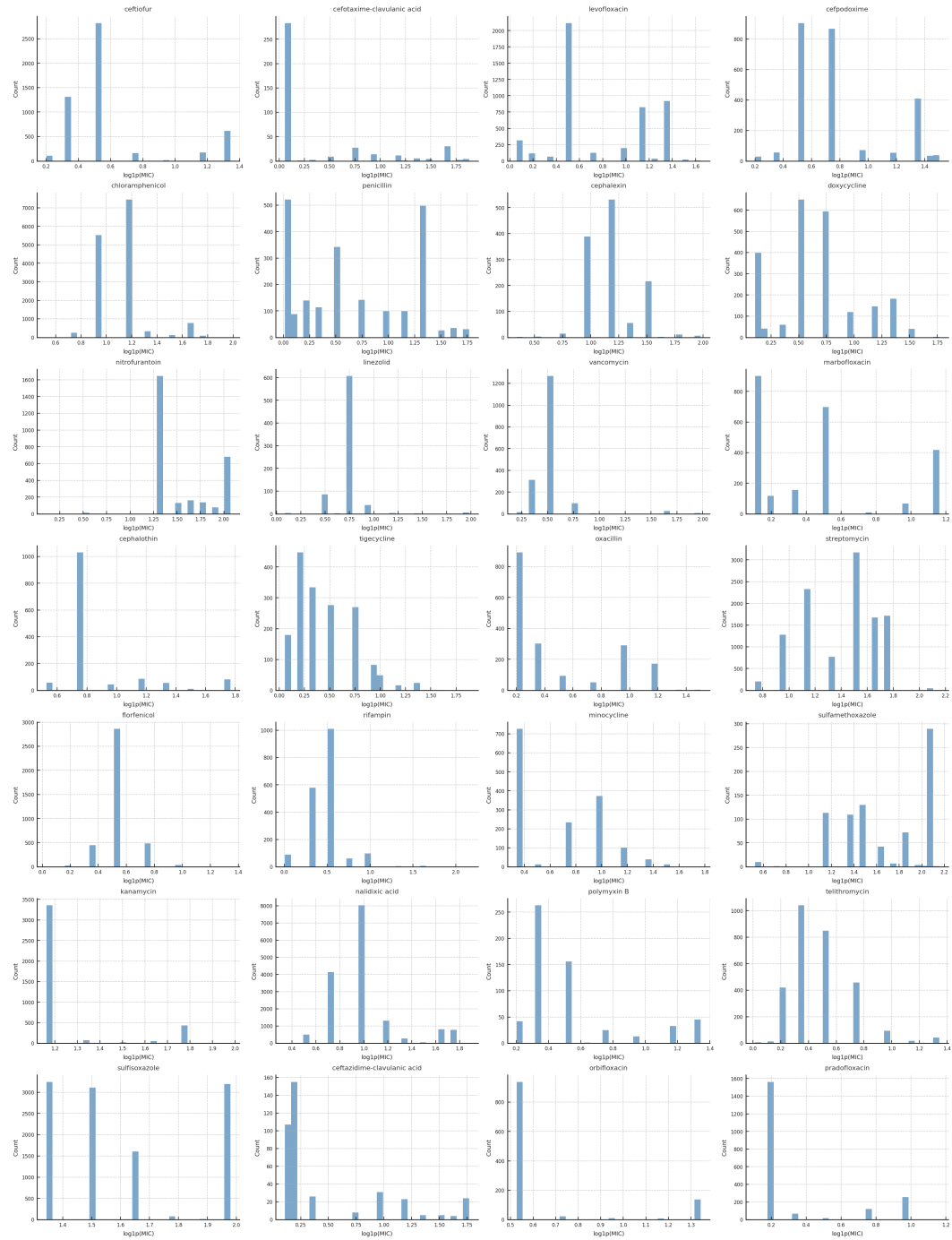
Regression. For regression MIC prediction, NCBI AST browser quotes minimum inhibitory concentrations (MICs) as $< x$, $\leq x$, $= x$, $\geq x$ and $> x$. These strings were translated to an actual number (y) by $y = x$ if NCBI quoted MIC= x , $\leq x$ or $\geq x$; $y = 2 \times x$ if MIC quoted as MIC $> x$, and $y = 0.5 \times x$ if MIC was $< x$. We filter out antibiotics with less than 500 available readings to ensure that the models are trained on a sufficiently large sample. This resulted in 56 antibiotics. To dampen the long tails, we normalized the MIC with a $\log_{1p}$ transformation. The final MIC distributions per antibiotic can be found in Supp. Fig 1 & 2.

Split. Due to low number of samples for many antibiotics and the variability between genomes, which may skew the results when using a single split, we trained and evaluated all models and antibiotics with k -fold split. Specifically, for each antibiotic we recommend: (1) splitting the available data into 5 equal splits using stratified split for the binary case and random split for regression. (2) In each split, further divide the larger train set into train and val, where validation makes up 20% of the train split. (3) Training the model on the train set and use the best performing model on the validation to evaluate the model on the test set.

Supplementary Table 2: Number of resistant, susceptible, and total labelled genomes for every antibiotic in the binary prediction setting.

Drug	# Resistant	# Susceptible	# Total
amikacin	7353	588	7941
ampicillin	10052	6211	16263
amoxicillin-clavulanic acid	12458	1683	14141
azithromycin	10497	2509	13006
cefazolin	1666	1897	3563
cefepime	3237	1188	4425
cefotaxime	566	446	1012
cefoxitin	12463	3902	16365
ceftazidime	2362	717	3079
ceftriaxone	13385	3469	16854
ciprofloxacin	17353	4280	21633
clindamycin	4327	337	4664
ertapenem	7708	373	8081
erythromycin	5605	1860	7465
fosfomycin	1186	396	1582
gentamicin	18242	2951	21193
imipenem	8485	409	8894
kanamycin	3768	410	4178
levofloxacin	3302	1106	4408
meropenem	11320	811	12131
nalidixic acid	1262	780	2042
nitrofurantoin	10535	5316	15851
oxacillin	6929	762	7691
piperacillin-tazobactam	3294	968	4262
rifampin	611	354	965
streptomycin	11965	2442	14407
sulfamethoxazole	3770	707	4477
sulfisoxazole	1569	323	1892
tetracycline	6936	3242	10178
tigecycline	662	357	1019
trimethoprim	2617	582	3199
ampicillin-sulbactam	1632	451	2083
aztreonam	2530	728	3258
ceftaroline	611	152	763
chloramphenicol	5345	1016	6361
colistin	572	290	862
daptomycin	492	131	623

Supplementary Figure 1: The $\log_{1p}(\text{MIC})$ distribution per antibiotic (1/2).

Supplementary Figure 2: The $\log_{1p}(\text{MIC})$ distribution per antibiotic (2/2).

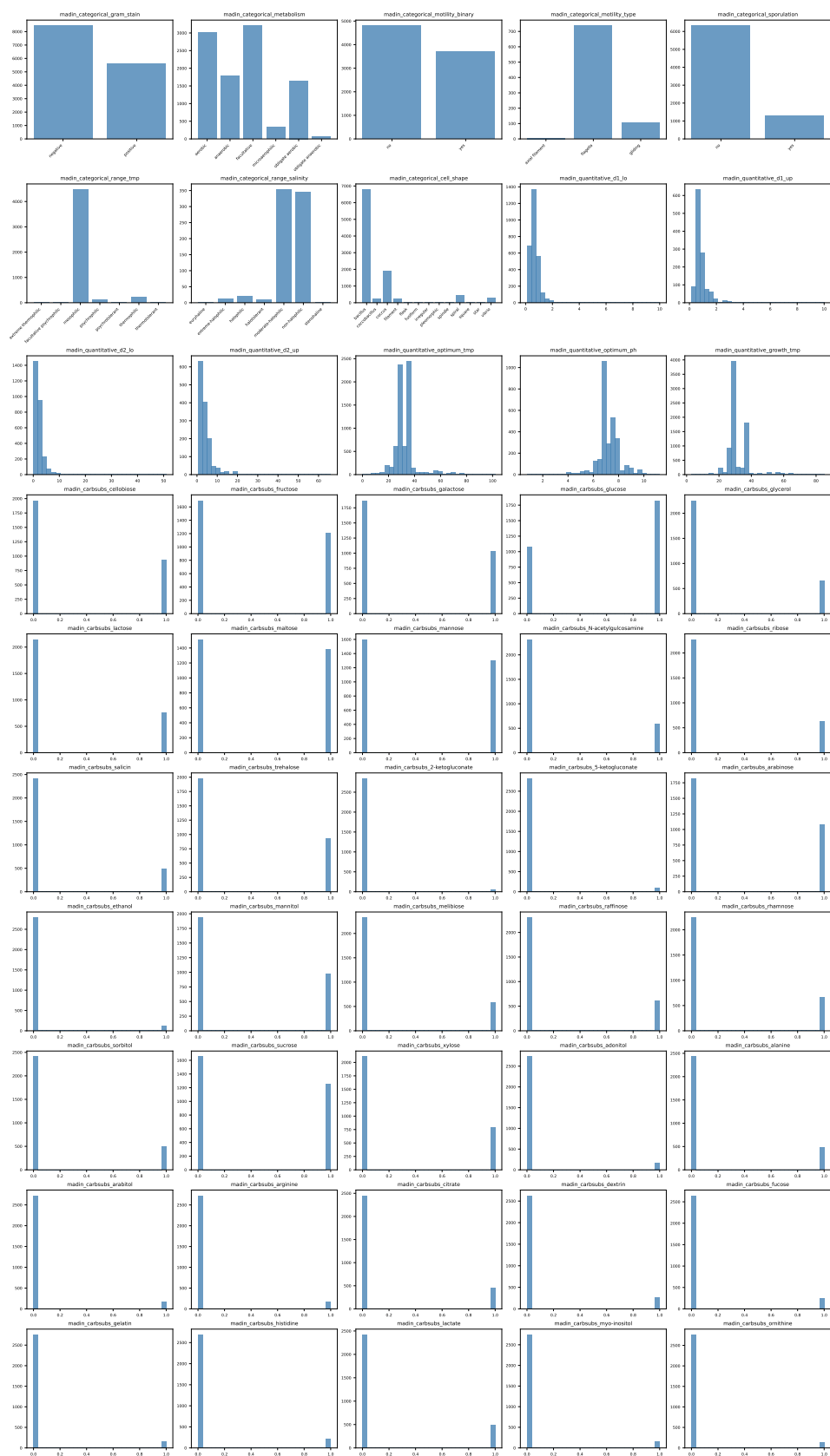
A.6 PHENOTYPIC TRAITS PREDICTION

We collected the phenotypic traits by collating two major sources (Madin et al., 2020; Weimann et al., 2016). To each phenotypic trait label we prepend its data source name so that the provenance of every label is explicit and any potential name collisions between the two catalogues are avoided. We keep duplicate traits that appear in both sources because they expand the number of labelled genomes without forcing us to merge measurements that were obtained with different experimental protocols. Using the taxonomy IDs and assembly accessions provided in the phenotypic traits datasets, we downloaded the associated genome DNA and protein sequences from the NCBI GenBank (<https://www.ncbi.nlm.nih.gov/genbank/>). We limit ourselves to categorical phenotypes and filter out the phenotypic traits with less than 500 genomes to ensure that the models are trained and evaluated on a sufficiently large sample. Additionally, we removed the classes with less than 50 samples. This resulted in 139 unique phenotypes across 24,462 genomes. We group the phenotypic traits into 5 distinct groups according to the type of biological information they capture (Supp. Fig. 3). We include the distributions of each phenotypic trait label in the Supp. Fig. 3-5 which shows large variation in the number of available labels per phenotypic trait which affects the final results.

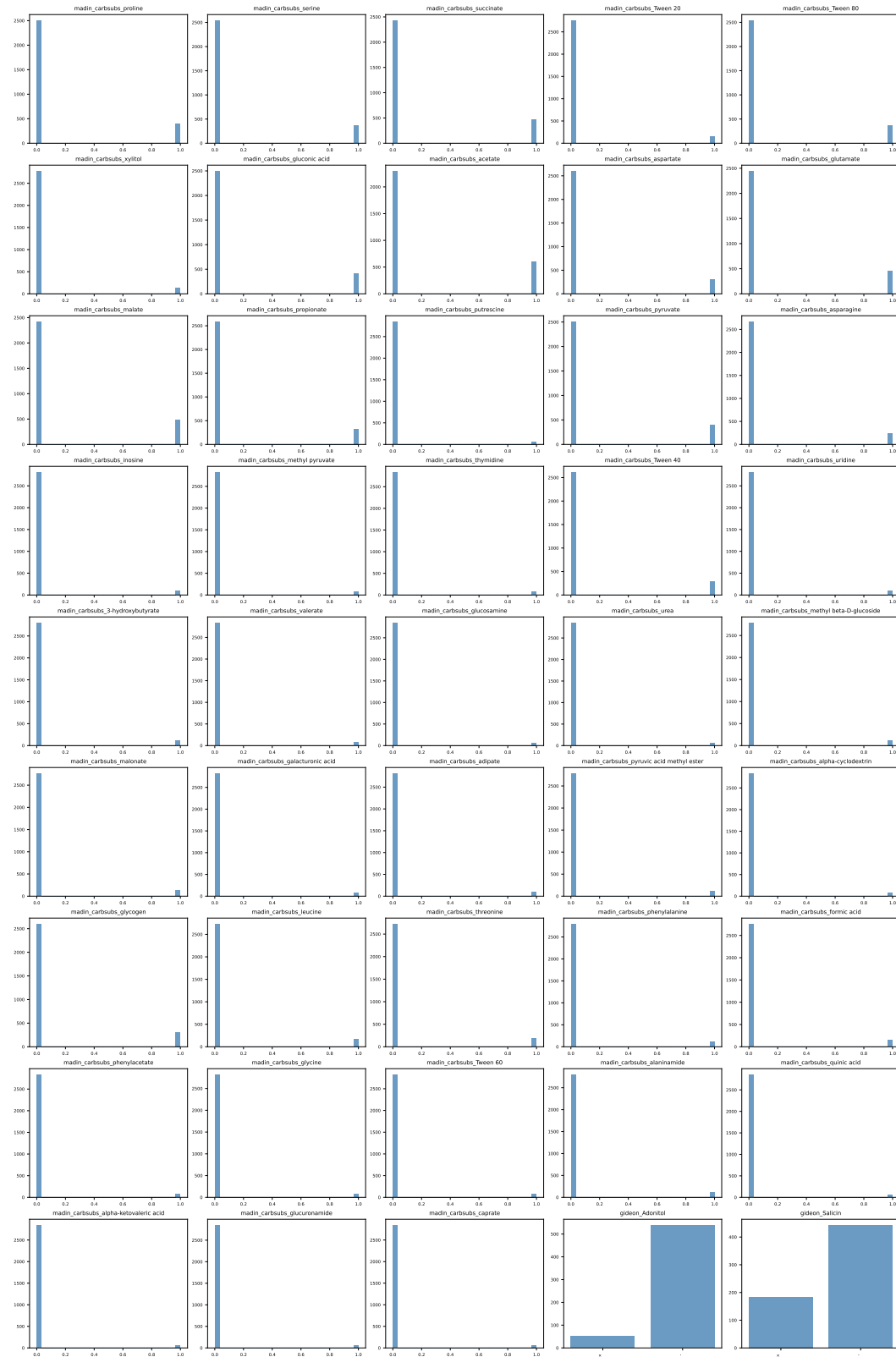
Split. For each phenotype, we split the data into 60/20/20 train, validation and test partitions respectively. As similar genomes often share the same phenotype, we split the data for each phenotype by genus—placing all genomes from a genus in one split—and evaluate on held-out genera, enforcing generalization to phylogenetically distant strains. To obtain stable estimates, as many traits are rare, we aggregate the per-phenotype results across 5 independent runs with different data splits.

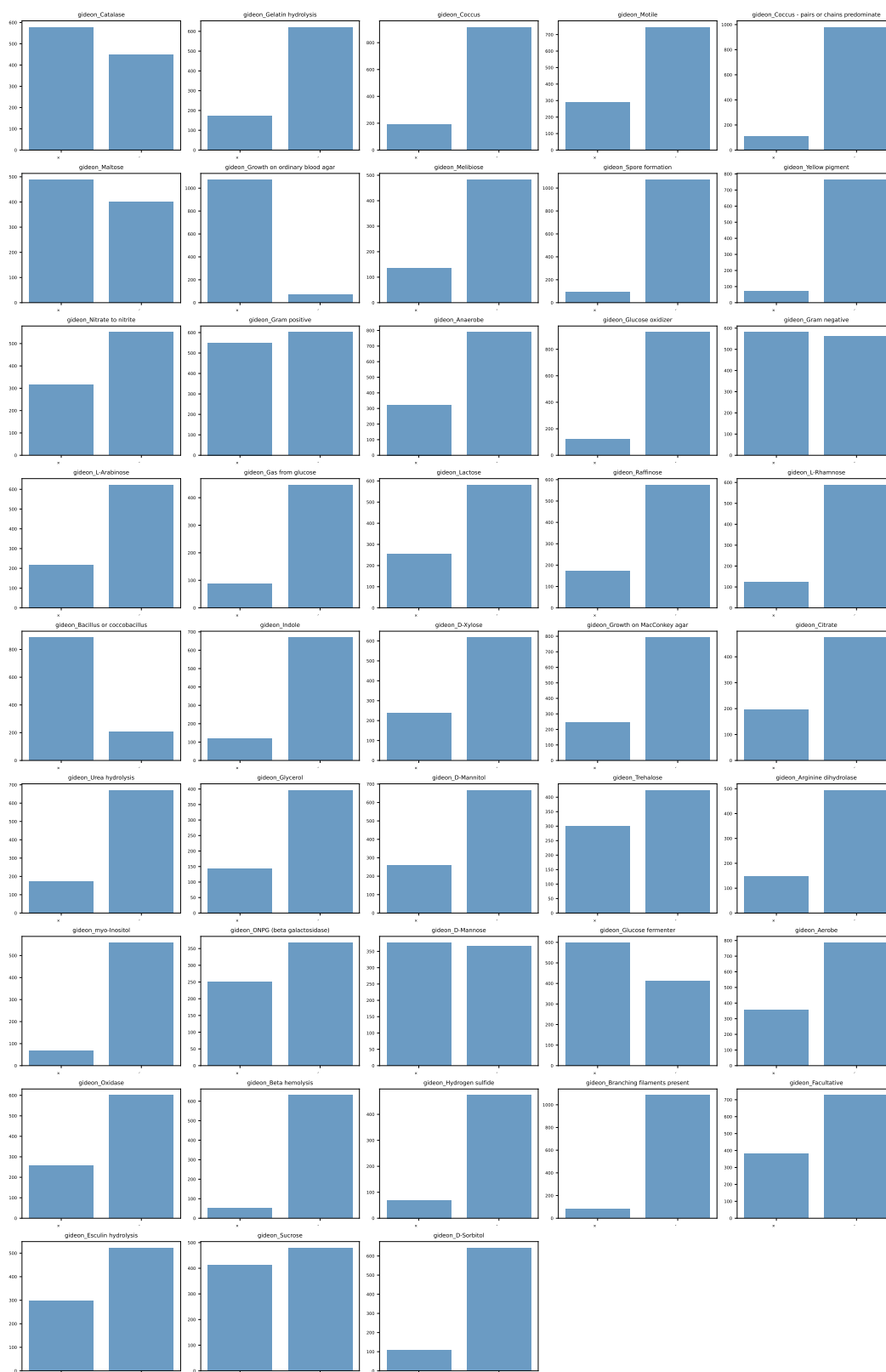
Supplementary Table 3: Phenotype groups and their associated phenotypes. The name before the first “_” symbolizes the dataset source. *madin* phenotypes were extracted from Madin et al. (2020) and *gideon* from Weimann et al. (2016).

Phenotype group	Phenotype
Biochemical Activity	gideon.Gelatin hydrolysis, gideon.Indole, gideon.Urea hydrolysis, gideon.Methyl red, gideon.VP (Voges Proskauer), gideon.-Gal (beta-galactosidase), gideon.Beta hemolysis, gideon.Hydrogen sulfide, gideon.Esculin hydrolysis
Carbon Utilization	madin.carbsubs_cellobiose, madin.carbsubs_glucose, madin.carbsubs_glycerol, madin.carbsubs_lactose, madin.carbsubs_maltose, madin.carbsubs_mannitol, madin.carbsubs_sucrose, madin.carbsubs_xylose, gideon.D-Arabitol, gideon.D-Mannose, gideon.Sucrose, gideon.D-Sorbitol
Growth Conditions	madin.categorical_range_tnp, madin.categorical_range_pH, madin.quantitative_optimum_tnp, madin.quantitative_optimum_pH, madin.quantitative_optimum_O2, madin.quantitative_growth_rate, gideon.Growth on ordinary blood agar, gideon.Growth on MacConkey agar
Morphology	madin.categorical_gram_stain, madin.categorical_sporulation, madin.categorical_cell_shape, madin.quantitative_average_cell_size, gideon.Motility, gideon.Spores, gideon.Shape: bacillus or coccobacillus, gideon.Branching filaments present
Respiration Metabolism	madin.categorical_metabolism, gideon.Nitrate reduction, gideon.Alanine aminopeptidase, gideon.O/F glucose oxidizer, gideon.Gas from glucose, gideon.Glucose fermenter, gideon.Aerobe, gideon.Oxidase, gideon.Facultative

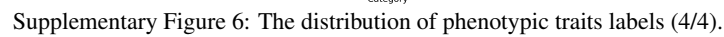


Supplementary Figure 3: The distribution of phenotypic traits labels (1/4).





Supplementary Figure 5: The distribution of phenotypic traits labels (3/4).



B ADDITIONAL MODELING & EVALUATION DETAILS

We outline the experimental details used for modeling & evaluation for all model types; DNA Language Models (DNA LMs), protein Language Models (pLMs) and bacterial Language Models (bLMs). Implementation and further details can be found at <https://anonymous.4open.science/r/BacBench-B6EF>. All of the modeling code has been implemented in PyTorch (Paszke, 2019).

DNA LMs. For every model we load the public checkpoint (Supp. Table 4) and keep the hidden states of the *last* encoder layer. If a DNA string has length of G tokens, the encoder produces $\mathbf{X} \in \mathbb{R}^{G \times D}$; we collapse it with a simple mean-pool to obtain a vector

$$\mathbf{z} = \frac{1}{G} \sum_{g=1}^G \mathbf{X}_{g,\cdot} \in \mathbb{R}^D.$$

Gene- & system-level. When the input gene (plus upstream promoter) is longer than the model context C , we slide a window of length C with stride s and embed each window. Averaging the M window vectors $\{\mathbf{z}^{(1)}, \dots, \mathbf{z}^{(M)}\}$ gives the final gene representation $\bar{\mathbf{z}}_{\text{gene}} = \frac{1}{M} \sum_m \mathbf{z}^{(m)}$.

Genome-level. We treat the whole genome in the same way: split into C -bp windows, embed each, and average; if several contigs are present we first average per-contig and then across contigs.

pLMs. A protein of K amino acids yields $\mathbf{X} \in \mathbb{R}^{K \times D}$ and $\mathbf{z} = \frac{1}{K} \sum_k \mathbf{X}_{k,\cdot}$. For genome-level tasks we aggregate the M protein vectors in that genome via $\bar{\mathbf{z}}_{\text{prot}} = \frac{1}{M} \sum_{i=1}^M \mathbf{z}_i$. Gene- and operon-level tasks use only the proteins involved.

bLMs. gLM2 ingests *mixed-modality* genomic scaffolds in which protein-coding regions are translated to amino-acid tokens and intergenic regions remain as nucleotide tokens. We tokenize each scaffold into a single sequence up to the model context C (4,096 tokens); longer scaffolds are processed with a sliding window of length C and stride s , mirroring the DNA LM setup. For each window, we retain the last-layer hidden states and mean-pool to obtain a window vector $\mathbf{z}^{(m)} \in \mathbb{R}^D$; genome-level embeddings average across windows (and across contigs when present):

$$\bar{\mathbf{z}}_{\text{genome}} = \frac{1}{M} \sum_{m=1}^M \mathbf{z}^{(m)}.$$

For gene essentiality and operon tasks, inputs include contextual flanking sequence to supply regulatory/positional cues, but the gene embedding is computed by averaging only the hidden states of tokens that fall within the gene’s coordinates; empirically this has shown to outperform averaging the entire slice.

Bacformer model takes as input *local* protein vectors $\mathbf{z}_1, \dots, \mathbf{z}_M$ obtained exactly as in the pLM setting. They are then **ordered by their genomic coordinates** (chromosome followed by plasmids) so that the model can “see” genome organisation. Rotary positional embeddings (Su et al., 2024) are added to these vectors and the ordered sequence is fed to a *genome-level* Transformer encoder (Vaswani et al., 2017) with L layers,

$$\mathbf{H}^{(0)} = [\mathbf{z}_1; \dots; \mathbf{z}_M], \quad \mathbf{H}^{(\ell+1)} = \text{Transformer}^{(\ell)}(\mathbf{H}^{(\ell)}), \quad \ell = 0, \dots, L-1.$$

The output $\mathbf{H}^{(L)} = [\tilde{\mathbf{z}}_1; \dots; \tilde{\mathbf{z}}_M]$ contains **contextualised protein embeddings** $\tilde{\mathbf{z}}_i \in \mathbb{R}^D$ that encode both the protein sequence and its genomic neighbourhood. During pre-training Bacformer learns to predict which proteins co-evolved, so these embeddings capture functional coupling across the genome.

Empirically, averaging token embeddings performed slightly better than using the special [CLS] token, so mean pooling was used for every model throughout the paper.

B.1 GENERAL TRAINING DETAILS.

For all tasks and settings which required finetuning except gene essentiality prediction, we kept the frozen backbone encoder model frozen, and only finetuned the neural network layer(s) stacked on

Supplementary Table 4: Summary of benchmarked models. “Max ctx.” = maximum context length supported at inference; for DNA models measured in base pairs, for pLMs in amino acids, and for Bacformer in number of proteins present in the genome. “dim” = dimensionality of the output of the last hidden layer. *We use either `bacformer-masked-complete-genomes` or `bacformer-masked-MAG` depending on the genome type used as input. The DNA LMs, pLMs and bLMs are separated by a horizontal line.

Model	Input	Variant / Checkpoint	Objective	Tokenisation	Params	dim	Training corpus	Max ctx.
Mistral-DNA (Mourad, 2025)	DNA	Mistral-DNA-v1-138M-bacteria	Autoregressive	Byte-pair	138M	768	Bacteria	512
DNABERT-2 (Zhou et al., 2023)	DNA	DNABERT-2-117M	Masked	Byte-pair	117M	768	Multi-kingdom	512
Nucleotide Transformer (Dalla-Torre et al., 2024)	DNA	nucleotide-transformer-v2-250m-multi-species	Masked	k-mer	250M	768	Multi-kingdom	2,048
ProkBERT (Ligeti et al., 2024)	DNA	neuralbioinfo/prokbert-mini-long	Masked	k-mer	27M	384	Bacteria	4,096
Evo (Nguyen et al., 2024)	DNA	evo-1-8k-base (1.1,fix)	Autoregressive	Single nucleotide	6.5B	4,096	Multi-kingdom	8,192
ESM-2 (Lin et al., 2022)	Single protein seq.	esm2_t12_35M_UR50D	Masked	Single amino acid	35M	480	Multi-kingdom	1,024
ESM-C (ESM Team, 2024)	Single protein seq.	esm_c_300m	Masked	Single amino acid	300M	960	Multi-kingdom	1,024
ProtBERT (Elmaghr et al., 2021)	Single protein seq.	prot_bert	Masked	Single amino acid	420M	1,024	Multi-kingdom	1,024
gLM2 (Corman et al., 2024)	Mixed modality (DNA & protein seq.)	tatatabio/glm2_650M	Masked	Single nucleotide/amino acid	650M	1,280	Bacteria	4,096
Bacformer (Wiatrak et al., 2025)	Multiple protein seq.	macwiatrak/bacformer-masked-complete-genomes*	Masked	Single protein	27M	480	Bacteria	6,000

top of the model encoder. This was motivated by the computational cost required to embed all 67k genomes with all models. Further details on runtime and computational cost can be found below. We used Adam optimizer (Kingma, 2014) in all finetuning setups. All of the checkpoints used have been downloaded directly from HuggingFace and are specified in Supp. Table 4. For each task and model combination, we tuned the learning rate keeping other parameters unchanged. Further details on hyperparameters used for each task and setup can be found in task-specific sections below.

Supplementary Table 5: Model-specific context parameters used in this study. “Max ctx.” is the maximum input length; “DNA-seq overlap” is the stride between consecutive windows when sliding across a genome; “Promoter length” is the upstream sequence length concatenated for promoter prediction.* For Bacformer the maximum input size of a protein is 1,024 amino acids and maximum number of proteins in the genome is set to 6,000.

Model	Max ctx.	DNA-seq overlap	Promoter length
Mistral-DNA	512	16	128
DNABERT-2	512	16	128
Nucleotide Transformer	2,048	32	128
ProkBERT	4,096	64	128
Evo	8,192	32	128
ProtBERT	1,024	N/A	N/A
ESM-2	1,024	N/A	N/A
ESM-C	1,024	N/A	N/A
gLM2	4,096	64	128
Bacformer	1,024 / 6,000*	N/A	N/A

B.2 TASK DETAILS

We outline the experimental details for each task, outlining the hyperparameters used and evaluation setup. All of the experiments have been performed on a single NVIDIA A100 with 32 CPU cores.

Gene essentiality prediction. We stacked a single linear layer preceded by a dropout of 0.2 and layer normalization (Ba et al., 2016) on top of the gene embeddings and trained it with binary cross-entropy loss to predict gene essentiality (1=essential, 0=non-essential). We trained the model for maximum of 100 epochs with early stopping patience of 10, monitoring the macro AUROC across genomes on the validation set. For each model we tuned the learning rate specified in Supp. Table 6. The weight decay for the Adam optimizer has been set to 0.02 for all models.

Evo. Evo natively does not provide straight-forward access to the output of the last hidden layer of the model. Therefore, we experimented with two ways of extracting the gene embeddings from Evo. Given a gene sequence G of size N , 1) we used the script provided as part of the Evo implementation (<https://github.com/evo-design/evo>) to score the log-likelihoods of the nucleotides in a sequence, resulting in a vector $z \in \mathbb{R}^N$, and 2) modified the Evo model to return the output of the last hidden state, resulting in a matrix $X \in \mathbb{R}^{N \times D}$, where D is model dimensionality which here equals 4,096. We then similarly as with other models took the average of all sequence tokens resulting in a vector $x \in \mathbb{R}^D$. We include this Evo implementation in the BacBench code repository (<https://anonymous.4open.science/r/BacBench-B6EF>). The option 1) yielded much better results on the validation set, therefore, we used it for final benchmarking.

Supplementary Table 6: Learning rates used for essential genes linear layer.

Method	Learning rate
Mistral-DNA	0.005
DNABERT-2	0.005
Nucleotide Transformer	0.005
ProkBERT	0.01
Evo	0.001
ProtBERT	0.005
ESM-2	0.005
ESM-C	0.005
gLM2	0.001
Bacformer	0.005

Operon identification. The operon-identification task is evaluated *zero-shot*; no fine-tuning is performed. We formulated operon prediction as a binary boundary classification problem, evaluating whether two contiguous genes form an operon. To address this, we developed a method that incorporates three features: (1) gene embeddings and (2) the strand of each gene.

First, we compute cosine similarity between adjacent genes using embeddings from the last hidden layer of pretrained models. We also record each gene’s transcriptional strand from the genome assembly.

We then define a simple score that combines similarity and strand co-orientation:

$$s_i = c_i I_{\text{strand}}, \quad c_i = \frac{1}{2}(1 + \cos(\hat{h}_i, \hat{h}_{i+1})),$$

where $I_{\text{strand}} = 1$ if both genes are on the same strand and 0 otherwise, and $\hat{h}_i, \hat{h}_{i+1}$ are the ℓ_2 -normalised embeddings of genes i and $i+1$. The score $s_i \in [0, 1]$ serves as the operon-membership score; a pair is classified as belonging to the same operon when s_i exceeds a threshold. Mapping cosine similarity to $[0, 1]$ stabilizes the scale across models, while the strand indicator provides a hard veto since genes on opposite strands do not belong to the same operon.

Performance is computed per strain.

Evo. Evo natively does not provide straight-forward access to the output of the last hidden layer of the model. Therefore, we experimented with two ways of extracting the gene embeddings from Evo. Given a gene sequence G of size N , 1) we used the script provided as part of the Evo implementation (<https://github.com/evo-design/evo>) to score the log-likelihoods of the nucleotides in a sequence, resulting in a vector $z \in \mathbb{R}^N$, and 2) modified the Evo model to return the output of the last hidden state, resulting in a matrix $X \in \mathbb{R}^{N \times D}$, where D is model dimensionality which here equals 4,096. We then similarly as with other models took the average of all sequence tokens resulting in a vector $x \in \mathbb{R}^D$. We include this Evo implementation in the BacBench code repository (<https://anonymous.4open.science/r/BacBench-B6EF>). The option 2) yielded significantly better results on operon identification, therefore, we used it for final benchmarking.

Protein-protein interaction prediction. To predict whether the two proteins interact, we fed the two protein embeddings into a linear model, which is trained to predict a binary label (1=interaction, 0=no interaction). The linear model is a single-layer neural network preceded by a dropout of 0.2 and layer normalization (Ba et al., 2016). The protein embeddings are fed into the linear classifier, after which the pairs of interacting proteins are averaged and passed through a final binary classification layer preceded by a dropout of 0.2. The model is trained to minimize the binary cross-entropy loss. We experimented with different learning rates for all the models and set the final learning rate to 0.001, which has shown to perform the best for all the models. We trained the model for the maximum epochs of 10 and no early stopping patience. We set the maximum gradient norm to 2.0 and monitor the validation loss across genomes. The weight decay for the Adam optimizer is set to 0.01.

Strain clustering. To compute the strain-clustering metrics we run Leiden clustering (Traag et al., 2019) over a grid of parameters. We vary the *resolution* in $[0.1, 0.25, 0.5, 1.0]$ and the

number of neighbours in $[5, 10, 15]$, evaluating every pairwise combination. Lower resolutions are omitted to avoid collapsing many genomes into a single (or just a few) giant clusters. After computing the clustering metrics for every parameter pair, we keep for each method the combination that maximises the mean performance across species-, genus- and family-level labels. The Leiden clustering is performed using the `scanpy` package (Wolf et al., 2018).

Antibiotic resistance prediction. To predict the antibiotic resistance, we train a linear model for each drug and method combination. The linear model is a single-layer neural network preceeded by a dropout of 0.2 and layer normalization (Ba et al., 2016). We train the models separately for the (i) binary and (ii) regression MIC prediction case. We optimize the former for the binary cross entropy loss and the latter for the mean squared error loss. We train all models for the maximum of 100 epochs with early stopping patience of 5, monitoring the validation AUPRC in the binary setup and validation R^2 in the regression setup. We experimented with various learning rates for each model, setting the final learning rate to 0.005 which attained the best results across folds and seeds for all models. We set the weight decay in the Adam optimizer to 0.01.

We have also experimented with training a multi-task linear model, which simultaneously predicts antibiotic resistance of a genome to multiple drugs, however, it performed worse than a separate linear model for each drug.

Phenotypic traits prediction. To predict a phenotypic trait from a genome-level embedding, we train an linear model for each phenotype and method combination. The linear model is a single-layer neural network preceeded by a dropout of 0.2 and layer normalization (Ba et al., 2016). As all labels are categorical, we optimize it to minimize the cross-entropy loss. We set the maximum number of epochs to 2,000 and early stopping patience to 50, monitoring the validation loss. The learning rate for all models was set to 0.01. We use the cross-entropy loss. To account for the class imbalance, we weigh each class according to:

$$w_c = \frac{N}{K n_c}, \quad (1)$$

where n_c is the number of training samples in class c , K is the total number of classes, and N is the total number of training samples.

Runtime analysis. A typical bacterial genome contains on the order of 3,000–5,000 genes and ~ 4 –6 Mbp of DNA, with average protein lengths of ~ 300 amino acids; embedding an entire genome therefore entails thousands of forward passes for protein LMs and millions of tokens for DNA LMs, making raw throughput a practical bottleneck for population-scale studies. We measured the wall-clock time required to embed the 11 genomes used in the operon identification task on a single NVIDIA A100 and extrapolated to BacBench’s full collection of 67,000 genomes (Table 7). The fastest models are **ESM-2** and **Bacformer** (64–67 s for 11 genomes; ≈ 1.2 – 1.25 k GPU-hours for 67k genomes). DNA LMs add a modest cost (e.g., 93 s for Mistral-DNA; 168 s for DNABERT-2) yet remain tractable on a single GPU. In stark contrast, the 6.5B-parameter **Evo** is $\sim 10^4\times$ slower (5.76×10^5 s for 11 genomes, i.e., ~ 14 h per genome), yielding an impractical $\sim 1.07 \times 10^7$ GPU-hours for 67k genomes on one GPU. These measurements underline that—even when accuracy is the primary goal—runtime quickly becomes the limiting factor for population-scale analyses.

Beyond embedding costs, fine-tuning on genome-level tasks requires backpropagating through *all* proteins per genome (median $\sim 2,506$ proteins across our datasets) or through multi-megabase DNA contexts, which dramatically amplifies memory and compute, requiring often ≥ 200 NVIDIA A100 GPUs for a single genome. Practically, this makes genome-scale fine-tuning out of reach for most academic labs for DNA LMs and pLMs, while remaining feasible for Bacformer-style bLMs; thus, linear-probe evaluations are a necessary, controlled proxy for model selection at scale.

B.3 DATASETS & MODELS LICENSES

To ensure that our datasets and benchmarks are reusable in the academic setting, we checked the license for each resource used in the manuscript. The Supp. Table 8 details a license for all models and datasets used in the manuscript.

Supplementary Table 7: Embedding runtime (s=seconds) on the operon identification task (11 genomes) on one NVIDIA A100 GPU and the extrapolated wall-clock time (h=hours) to process the entire collection of 67k bacterial genomes on the same hardware.

Model	Runtime (s)	Estimated time for 67 k genomes (h)
Mistral-DNA	93	1,731
DNABERT-2	168	3,127
Nucleotide Transformer	100	1,861
ProkBERT	106	1,973
Evo	575,629	10,712,540
ProtBERT	95	1,768
ESM-2	64	1,191
ESM-C	137	2,550
gLM2	182	3,387
Bacformer	67	1,247

Supplementary Table 8: External resources used in this study and their licences.

Resource	Type	Licence
ESM-2	Model	MIT
ESM-C	Model	Cambrian Open License Agreement
ProtBERT	Model	Academic Free License v. 3.0
gLM2	Model	Apache 2.0
Bacformer	Model	Apache 2.0
DNABERT-2	Model	Apache 2.0
ProkBERT	Model	MIT
Evo	Model	Apache 2.0
Nucleotide Transformer	Model	CC BY-NC-SA 4.0
Mistral-DNA	Model	Apache 2.0
MGnify	Data	CC0 1.0 Universal
NCBI AST Browser	Data	Public domain (U.S. Gov data)
NCBI GenBank/RefSeq	Data	Public domain (U.S. Gov data)
Database of Essential Genes	Data	CC BY 4.0
Operon DB	Data	CC BY 4.0
STRING DB	Data	CC BY 4.0

C ADDITIONAL RESULTS

We show and discuss further results across all tasks. To allow for comparability in future work, we include tables with all numerical results as well as per antibiotic and phenotypic traits scores.

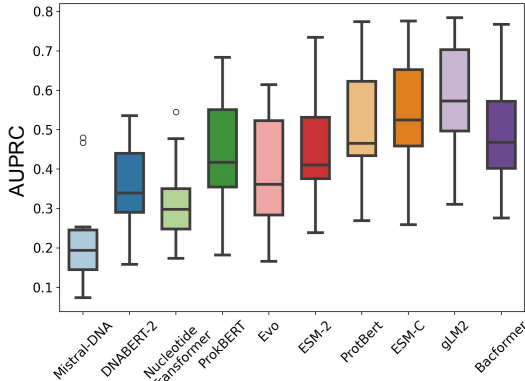
C.1 GENE ESSENTIALITY PREDICTION

The results on AUPRC show the bLMs and pLMs outperforming DNA LMs (Supp. Fig 7). gLM2 performs the best, showing the benefits of taking as input both DNA as well as protein sequence information. Moreover, increasing the pLM size appears to further boost performance, as demonstrated by ESM-C (300M) and ProtBERT (420M) outperforming ESM-2 (35M). Bacformer outperforms its protein representation backbone ESM-2, showing the benefits of incorporating whole-genome context. Evo outperforms other DNA LMs, except ProkBERT, demonstrating the performance gain by conducting pretraining on a relevant corpus. Finally the results on AUPRC show that there is large room for improvement. We believe that increasing the number and diversity of annotated genomes would significantly boost model performance.

The Supp. Table 9 shows the exact results across AUROC and AUPRC measured across distinct genomes.

Finetuning performance. To further analyse model performance, we have conducted finetuning on the gene essentiality task. The results show that finetuning boosts performance (Table 10); however, the model ranking remains largely unchanged, with the pLMs and bLMs outperforming DNA LMs. We have excluded Evo from finetuning due to the computational complexity required to finetune

Evo (Table 7). Gene essentiality is the only task for which we finetuned all models; genome-level tasks such as antibiotic-resistance prediction require end-to-end, whole-genome context and are prohibitively expensive for all models except Bacformer—underscoring the need for models that can be finetuned efficiently for whole-genome tasks.



Supplementary Figure 7: AUPRC across test genomes on the gene-essentiality prediction task. The box spans the inter-quartile range with a line marking the median value.

Supplementary Table 9: Overall performance on gene-essentiality prediction. Values are *mean* $\pm$ *standard deviation* over 3 random seeds; the best score for each metric is highlighted in bold.

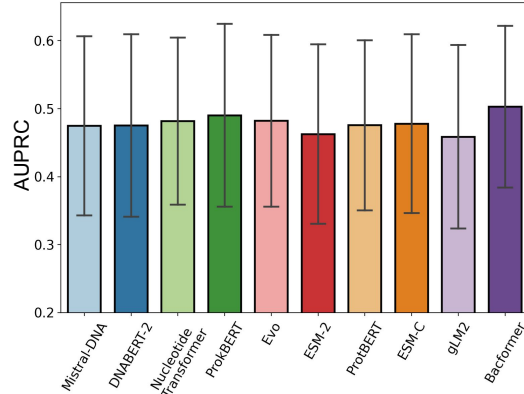
Method	AUROC	AUPRC
Mistral-DNA	57.52 $\pm$ 2.86	22.69 $\pm$ 14.21
DNABERT-2	68.21 $\pm$ 7.17	35.76 $\pm$ 12.46
Nucleotide Transformer	67.03 $\pm$ 5.56	31.78 $\pm$ 11.88
ProkBERT	74.79 $\pm$ 6.22	44.79 $\pm$ 15.25
Evo	73.71 $\pm$ 7.10	39.08 $\pm$ 15.21
ESM-2	77.99 $\pm$ 5.82	46.23 $\pm$ 15.87
ESM-C	82.25 $\pm$ 6.04	55.08 $\pm$ 15.83
ProtBERT	80.72 $\pm$ 5.85	52.00 $\pm$ 15.90
gLM2	83.77 $\pm$ 6.17	58.61 $\pm$ 14.96
Bacformer	80.72 $\pm$ 5.87	50.33 $\pm$ 15.43

Supplementary Table 10: Overall finetuning performance on gene-essentiality prediction. Values are *mean* $\pm$ *standard deviation* over 3 random seeds; the best score for each metric is highlighted in bold.

Method	AUROC	AUPRC
Mistral-DNA	62.18 $\pm$ 7.87	28.81 $\pm$ 13.51
DNABERT-2	74.88 $\pm$ 16.75	56.12 $\pm$ 29.71
Nucleotide Transformer	73.06 $\pm$ 14.99	48.56 $\pm$ 28.65
ProkBERT	69.88 $\pm$ 8.60	44.98 $\pm$ 14.36
ESM-2	85.36 $\pm$ 9.12	64.42 $\pm$ 17.79
ESM-C	89.96 $\pm$ 7.53	71.85 $\pm$ 14.49
ProtBERT	91.31 $\pm$ 73.02	73.02 $\pm$ 16.56
gLM2	90.12 $\pm$ 7.87	74.03 $\pm$ 21.70
Bacformer	90.31 $\pm$ 7.40	72.13 $\pm$ 15.16

C.2 OPERON IDENTIFICATION

The AUPRC follows the same trend as the AUROC described in the main manuscript: Bacformer attains the best scores, ProkBERT ranks second, and the remaining DNA LMs and pLMs form a tight cluster with broadly comparable performance—reinforcing that genome-level context and bacteria-specific pretraining matter more than modality alone; absolute AUPRC values are lower (as expected under class imbalance) but track the same model ordering.



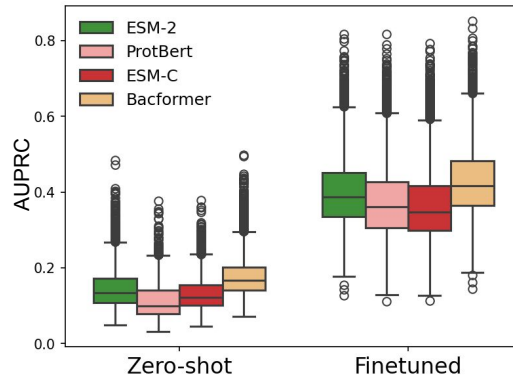
Supplementary Figure 8: AUPRC across test genomes on gene essentiality prediction task. The box spans the inter-quartile range with a line marking the median value.

Supplementary Table 11: Overall performance on operon identification. Values are *mean $\pm$ standard deviation*; the best score for each metric is highlighted in bold.

Method	AUROC	AUPRC
Mistral-DNA	74.16 $\pm$ 9.70	47.48 $\pm$ 29.47
DNABERT-2	74.13 $\pm$ 10.14	47.53 $\pm$ 30.07
Nucleotide Transformer	74.50 $\pm$ 8.57	48.19 $\pm$ 27.47
ProkBERT	75.77 $\pm$ 9.69	49.03 $\pm$ 30.07
Evo	74.77 $\pm$ 8.33	48.21 $\pm$ 28.27
ESM-2	73.02 $\pm$ 10.30	46.27 $\pm$ 29.55
ESM-C	75.25 $\pm$ 9.26	47.79 $\pm$ 29.39
ProtBERT	75.11 $\pm$ 8.65	47.58 $\pm$ 27.98
gLM2	72.85 $\pm$ 10.77	45.86 $\pm$ 30.18
Bacformer	77.59 $\pm$ 6.64	50.31 $\pm$ 26.61

C.3 PROTEIN-PROTEIN INTERACTION PREDICTION

The protein-protein interaction (PPI) prediction results on AUPRC (Supp. Fig. 9) show that contextual pLM, Bacformer, consistently outperforms other methods. We credit it to its usage of the genomic context. Moreover, the performance does not increase by scaling the model size, as shown by ESM-C and ProtBERT underperforming ESM-2. The overall results (Supp. Table 12) show relatively low performance considering the trainins set size ($\approx 7k$ genomes) even in the finetuned setting. We believe this is due to the 1) complexity of the task, 2) noisy source data, as STRING DB where the interactions have been extracted from collates information from a variety of sources and is not limited to experimentally validated interactions, highlighting the importance of building high quality PPI datasets.



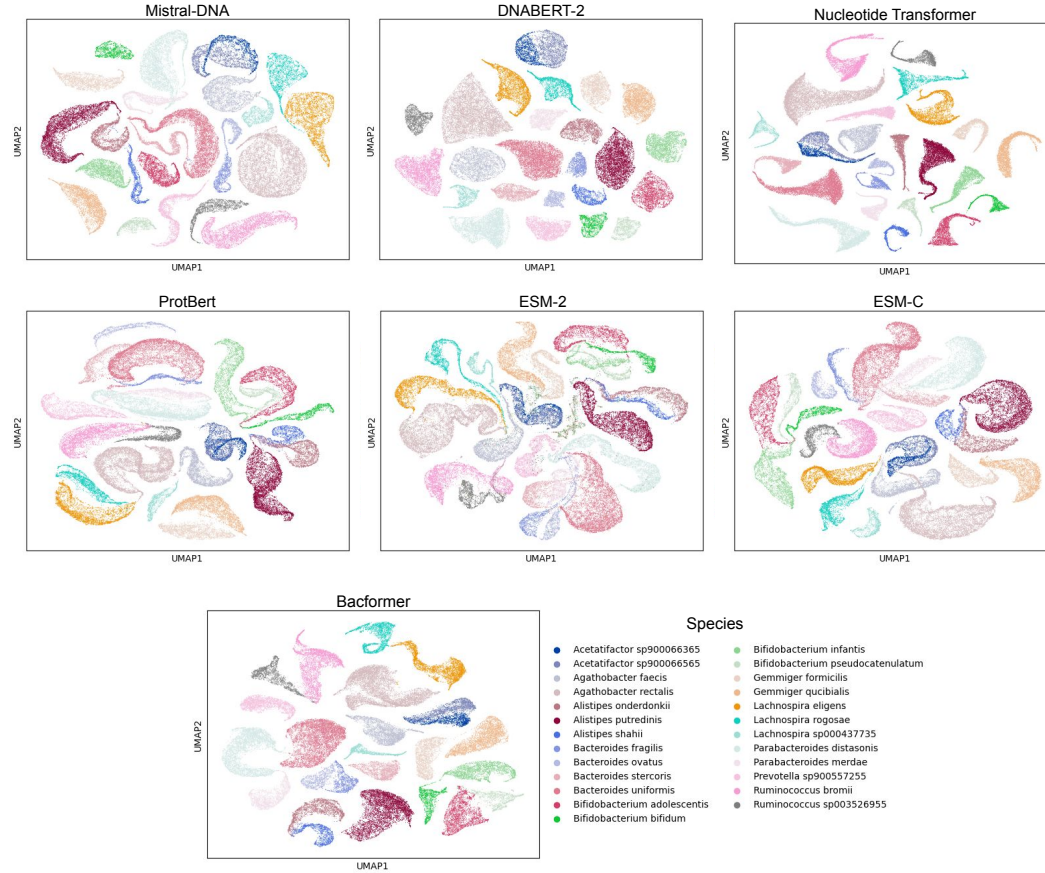
Supplementary Figure 9: AUPRC across test genomes on the protein-protein-interaction task in both zero-shot and fine-tuned settings. The box spans the inter-quartile range with a line marking the median.

Supplementary Table 12: Protein–protein–interaction prediction on the held-out genomes. Values are *mean $\pm$ standard deviation* over five seeds; the best score for each metric is shown in bold.

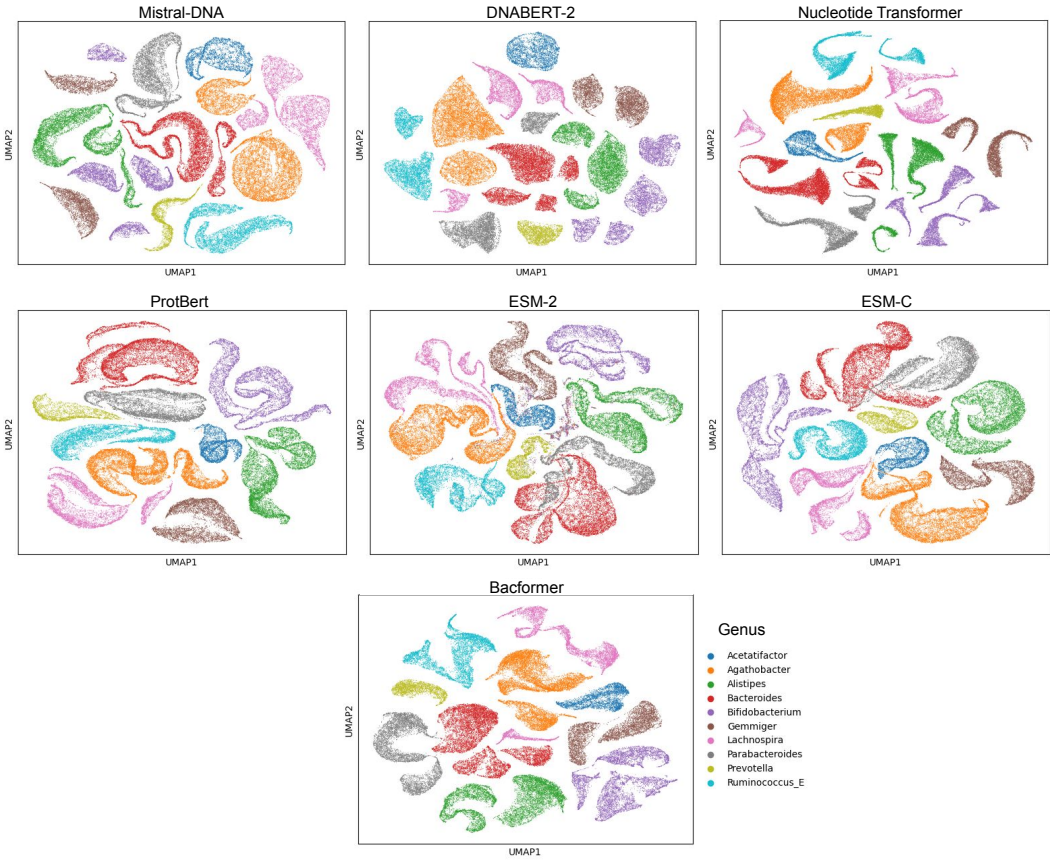
Method	Zero-shot		Finetuned	
	AUROC	AUPRC	AUROC	AUPRC
ESM-2	56.46 $\pm$ 2.10	14.94 $\pm$ 6.04	76.62 $\pm$ 2.35	40.68 $\pm$ 10.71
ProtBERT	47.20 $\pm$ 3.88	11.45 $\pm$ 4.89	74.05 $\pm$ 2.53	38.15 $\pm$ 11.16
ESM-C	55.17 $\pm$ 2.66	13.33 $\pm$ 4.61	74.21 $\pm$ 2.37	37.30 $\pm$ 10.98
Bacformer	63.09 $\pm$ 2.73	18.20 $\pm$ 6.28	79.09 $\pm$ 2.25	43.47 $\pm$ 10.61

C.4 STRAIN CLUSTERING

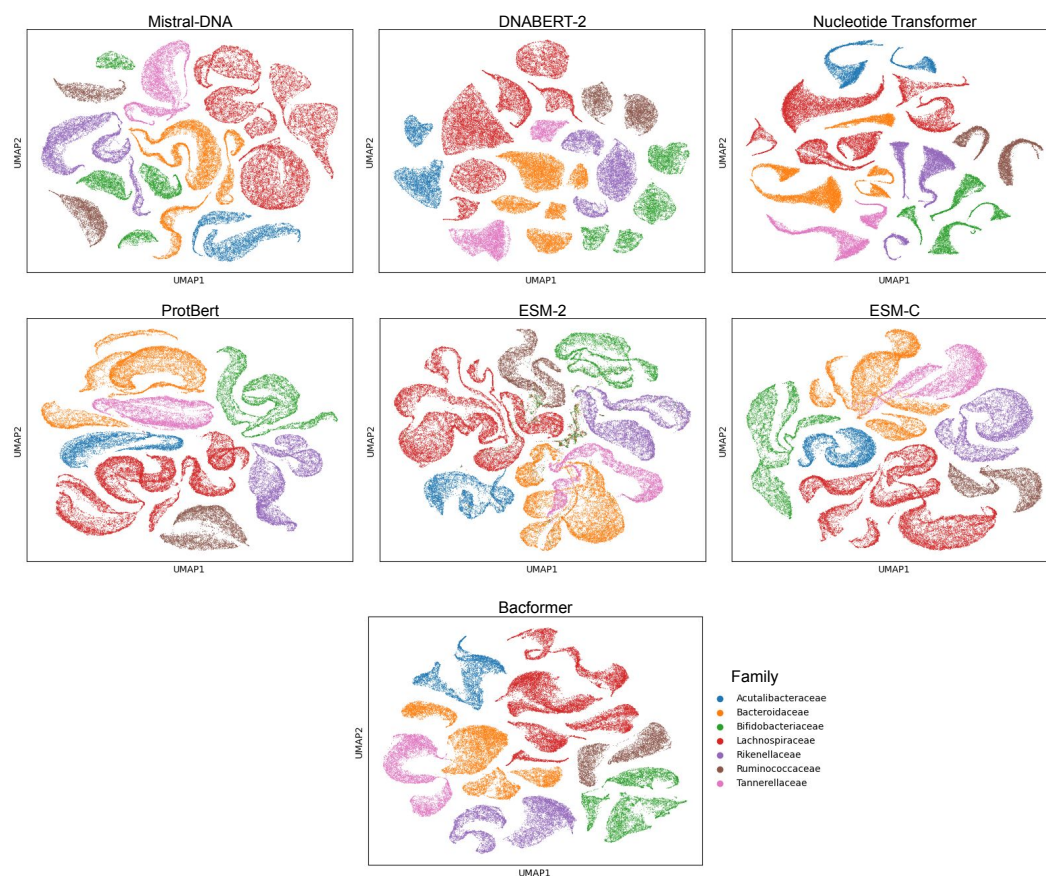
In addition to the strain clustering metrics included in the main manuscript, we plotted 2-dimensional UMAP results colored by species (Supp. Fig. 10), genus (Supp. Fig. 11) and family (Supp. Fig. 12) for a subset of models for further investigation. The UMAPs show that the strain representations differ between the models. All models cluster strains into separate species clusters, however, not all of them retain the phylogenetic similarities between species in the same genus or family. We also notice that the DNA LMs tend to output less overlapping species clusters, which boosts its performance at species level, but leads to lower results at higher taxonomic levels (genus and family).



Supplementary Figure 10: UMAP plots of strains (i.e. genome-level embeddings) across diverse methods colored by species.



Supplementary Figure 11: UMAP plots of strains (i.e. genome-level embeddings) across diverse colored by genus.



Supplementary Figure 12: UMAP plots of strains (i.e. genome-level embeddings) across diverse methods colored by family.

C.5 ANTIBIOTIC RESISTANCE PREDICTION

On the following pages we present the results per antibiotic across metrics. This includes binary prediction results (susceptible/resistant, Supp. Table 13) and MIC regression prediction results (Supp. Table 14). Each antibiotic was run across 5-folds with 3 random seeds to avoid variance stemming from random initialization and data-split bias. In the binary prediction setting the bLM-Bacformer outperforms other methods, achieving the best AUROC on 26 drugs, AUPRC on 27 drugs and F1 on 24 drugs out of 37 in total. Thus, showcasing the benefits of considering the interactions between proteins present in the genome. In the regression setting, Bacformer attains the best result on 44 drugs on R^2 and 45 on *Pearson* correlation coefficient out of total of 56 antibiotics. Finally, we see large variance across as well as within drugs. The former can be partly explained by the variable number of labels available per antibiotic, while the latter shows that, even within a single antibiotic, model performance can fluctuate markedly across folds and random seeds—highlighting sensitivity to sample composition and pointing to the need for larger, more balanced datasets or stronger regularisation to obtain stabler estimates.

Supplementary Table 13: Per-antibiotic binary antibiotic resistance prediction. Values are mean $\pm$ standard deviation across 3 seeds. Bold indicates the highest mean for each metric.

Antibiotic	Mistral-DNA		DNABERT2		Nucleotide Transformer		ProkBERT		ESM-2		ESM-C		ProtBERT		gLM2		Bioformer	
	AUROC	AUPRC	AUROC	AUPRC	AUROC	AUPRC	AUROC	AUPRC	AUROC	AUPRC	AUROC	AUPRC	AUROC	AUPRC	AUROC	AUPRC	AUROC	AUPRC
amikacin	90.24 $\pm$ 1.09	33.09 $\pm$ 1.79	88.46 $\pm$ 6.35	33.23 $\pm$ 9.22	90.77 $\pm$ 1.07	34.84 $\pm$ 3.83	90.47 $\pm$ 0.46	30.44 $\pm$ 0.30	91.62 $\pm$ 0.47	41.33 $\pm$ 0.75	91.94 $\pm$ 0.94	48.04 $\pm$ 3.46	91.70 $\pm$ 2.55	49.59 $\pm$ 6.31	88.94 $\pm$ 1.87	40.32 $\pm$ 2.48	94.43 $\pm$ 1.12	59.00 $\pm$ 4.79
amoxicillin-clavulanic acid	69.69 $\pm$ 1.80	28.06 $\pm$ 1.91	76.07 $\pm$ 1.35	37.21 $\pm$ 9.22	76.99 $\pm$ 2.30	41.06 $\pm$ 3.83	78.61 $\pm$ 0.50	40.47 $\pm$ 1.64	82.51 $\pm$ 0.67	48.70 $\pm$ 3.51	85.17 $\pm$ 1.03	54.77 $\pm$ 2.83	80.40 $\pm$ 1.07	48.71 $\pm$ 4.38	78.33 $\pm$ 1.48	40.72 $\pm$ 3.39	94.38 $\pm$ 1.54	59.03 $\pm$ 4.79
ampicillin	74.49 $\pm$ 1.42	68.27 $\pm$ 1.68	76.99 $\pm$ 1.32	69.81 $\pm$ 2.23	78.94 $\pm$ 0.94	72.00 $\pm$ 1.70	78.11 $\pm$ 0.47	91.77 $\pm$ 1.00	79.42 $\pm$ 0.24	73.05 $\pm$ 0.70	82.67 $\pm$ 0.59	77.14 $\pm$ 0.88	80.28 $\pm$ 0.47	45.55 $\pm$ 0.35	78.97 $\pm$ 0.67	71.79 $\pm$ 1.40	82.72 $\pm$ 0.84	78.85 $\pm$ 0.98
ampicillin-sulbactam	73.61 $\pm$ 1.84	68.27 $\pm$ 1.68	76.99 $\pm$ 1.32	69.81 $\pm$ 2.23	78.94 $\pm$ 0.94	72.00 $\pm$ 1.70	78.11 $\pm$ 0.47	91.77 $\pm$ 1.00	79.42 $\pm$ 0.24	73.05 $\pm$ 0.70	82.67 $\pm$ 0.59	77.14 $\pm$ 0.88	80.28 $\pm$ 0.47	45.55 $\pm$ 0.35	78.97 $\pm$ 0.67	71.79 $\pm$ 1.40	82.72 $\pm$ 0.84	78.85 $\pm$ 0.98
aztreonam	73.61 $\pm$ 1.84	68.27 $\pm$ 1.68	76.99 $\pm$ 1.32	69.81 $\pm$ 2.23	78.94 $\pm$ 0.94	72.00 $\pm$ 1.70	78.11 $\pm$ 0.47	91.77 $\pm$ 1.00	79.42 $\pm$ 0.24	73.05 $\pm$ 0.70	82.67 $\pm$ 0.59	77.14 $\pm$ 0.88	80.28 $\pm$ 0.47	45.55 $\pm$ 0.35	78.97 $\pm$ 0.67	71.79 $\pm$ 1.40	82.72 $\pm$ 0.84	78.85 $\pm$ 0.98
aztreonam	90.11 $\pm$ 0.68	90.06 $\pm$ 0.97	88.46 $\pm$ 1.67	89.37 $\pm$ 1.53	93.44 $\pm$ 0.90	93.24 $\pm$ 0.84	93.89 $\pm$ 1.14	92.53 $\pm$ 1.09	95.30 $\pm$ 0.39	94.25 $\pm$ 0.74	95.76 $\pm$ 0.37	94.74 $\pm$ 0.56	93.74 $\pm$ 0.81	93.40 $\pm$ 0.99	92.54 $\pm$ 0.65	92.28 $\pm$ 0.34	93.91 $\pm$ 1.21	93.09 $\pm$ 0.90
cefazolin	78.99 $\pm$ 1.58	83.03 $\pm$ 0.82	81.08 $\pm$ 1.32	84.62 $\pm$ 0.78	85.30 $\pm$ 2.61	89.63 $\pm$ 2.47	85.18 $\pm$ 2.72	90.00 $\pm$ 2.71	88.81 $\pm$ 3.86	91.90 $\pm$ 2.71	92.23 $\pm$ 1.25	94.13 $\pm$ 1.11	86.01 $\pm$ 0.51	90.29 $\pm$ 0.50	86.99 $\pm$ 0.26	90.97 $\pm$ 0.34	93.61 $\pm$ 0.86	90.75 $\pm$ 0.86
cefepime	71.17 $\pm$ 3.51	42.30 $\pm$ 3.73	71.83 $\pm$ 4.55	44.05 $\pm$ 4.25	75.30 $\pm$ 2.94	51.49 $\pm$ 6.06	76.21 $\pm$ 2.79	52.50 $\pm$ 2.15	79.50 $\pm$ 2.92	57.81 $\pm$ 12.81	82.08 $\pm$ 2.32	65.81 $\pm$ 4.89	78.94 $\pm$ 3.77	63.77 $\pm$ 1.15	76.11 $\pm$ 1.09	52.36 $\pm$ 2.19	84.75 $\pm$ 2.00	68.73 $\pm$ 5.33
cefotaxime	56.87 $\pm$ 1.73	79.94 $\pm$ 0.25	66.74 $\pm$ 3.36	89.85 $\pm$ 2.20	84.73 $\pm$ 3.23	95.39 $\pm$ 1.35	84.76 $\pm$ 4.79	94.30 $\pm$ 2.69	91.83 $\pm$ 2.39	97.54 $\pm$ 1.26	92.32 $\pm$ 2.25	97.78 $\pm$ 1.05	90.98 $\pm$ 2.58	97.36 $\pm$ 1.55	84.63 $\pm$ 2.64	95.96 $\pm$ 0.69	91.87 $\pm$ 2.58	97.83 $\pm$ 0.89
cefotaxime	76.41 $\pm$ 1.43	53.21 $\pm$ 0.34	80.23 $\pm$ 3.29	56.91 $\pm$ 4.44	82.94 $\pm$ 2.25	64.42 $\pm$ 2.09	82.10 $\pm$ 1.94	60.50 $\pm$ 1.73	85.36 $\pm$ 1.20	65.00 $\pm$ 1.26	83.56 $\pm$ 0.70	68.87 $\pm$ 1.68	84.79 $\pm$ 2.05	63.05 $\pm$ 5.57	81.91 $\pm$ 1.59	61.63 $\pm$ 1.07	80.03 $\pm$ 0.97	72.49 $\pm$ 2.16
ceftriaxone	74.81 $\pm$ 3.52	77.04 $\pm$ 3.62	81.13 $\pm$ 1.98	55.63 $\pm$ 4.02	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	87.57 $\pm$ 3.88	68.11 $\pm$ 0.96	83.59 $\pm$ 0.71	52.22 $\pm$ 0.87	81.33 $\pm$ 3.14	43.22 $\pm$ 0.21	77.05 $\pm$ 3.83	30.71 $\pm$ 4.56	88.26 $\pm$ 0.60	88.26 $\pm$ 0.60
ceftriaxone	74.81 $\pm$ 3.52	77.04 $\pm$ 3.62	81.13 $\pm$ 1.98	55.63 $\pm$ 4.02	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	87.57 $\pm$ 3.88	68.11 $\pm$ 0.96	83.59 $\pm$ 0.71	52.22 $\pm$ 0.87	81.33 $\pm$ 3.14	43.22 $\pm$ 0.21	77.05 $\pm$ 3.83	30.71 $\pm$ 4.56	88.26 $\pm$ 0.60	88.26 $\pm$ 0.60
ceftriaxone	74.81 $\pm$ 3.52	77.04 $\pm$ 3.62	81.13 $\pm$ 1.98	55.63 $\pm$ 4.02	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	87.57 $\pm$ 3.88	68.11 $\pm$ 0.96	83.59 $\pm$ 0.71	52.22 $\pm$ 0.87	81.33 $\pm$ 3.14	43.22 $\pm$ 0.21	77.05 $\pm$ 3.83	30.71 $\pm$ 4.56	88.26 $\pm$ 0.60	88.26 $\pm$ 0.60
ceftriaxone	74.81 $\pm$ 3.52	77.04 $\pm$ 3.62	81.13 $\pm$ 1.98	55.63 $\pm$ 4.02	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	87.57 $\pm$ 3.88	68.11 $\pm$ 0.96	83.59 $\pm$ 0.71	52.22 $\pm$ 0.87	81.33 $\pm$ 3.14	43.22 $\pm$ 0.21	77.05 $\pm$ 3.83	30.71 $\pm$ 4.56	88.26 $\pm$ 0.60	88.26 $\pm$ 0.60
ceftriaxone	74.81 $\pm$ 3.52	77.04 $\pm$ 3.62	81.13 $\pm$ 1.98	55.63 $\pm$ 4.02	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	87.57 $\pm$ 3.88	68.11 $\pm$ 0.96	83.59 $\pm$ 0.71	52.22 $\pm$ 0.87	81.33 $\pm$ 3.14	43.22 $\pm$ 0.21	77.05 $\pm$ 3.83	30.71 $\pm$ 4.56	88.26 $\pm$ 0.60	88.26 $\pm$ 0.60
ceftriaxone	74.81 $\pm$ 3.52	77.04 $\pm$ 3.62	81.13 $\pm$ 1.98	55.63 $\pm$ 4.02	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	87.57 $\pm$ 3.88	68.11 $\pm$ 0.96	83.59 $\pm$ 0.71	52.22 $\pm$ 0.87	81.33 $\pm$ 3.14	43.22 $\pm$ 0.21	77.05 $\pm$ 3.83	30.71 $\pm$ 4.56	88.26 $\pm$ 0.60	88.26 $\pm$ 0.60
ceftriaxone	74.81 $\pm$ 3.52	77.04 $\pm$ 3.62	81.13 $\pm$ 1.98	55.63 $\pm$ 4.02	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	87.57 $\pm$ 3.88	68.11 $\pm$ 0.96	83.59 $\pm$ 0.71	52.22 $\pm$ 0.87	81.33 $\pm$ 3.14	43.22 $\pm$ 0.21	77.05 $\pm$ 3.83	30.71 $\pm$ 4.56	88.26 $\pm$ 0.60	88.26 $\pm$ 0.60
ceftriaxone	74.81 $\pm$ 3.52	77.04 $\pm$ 3.62	81.13 $\pm$ 1.98	55.63 $\pm$ 4.02	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	87.57 $\pm$ 3.88	68.11 $\pm$ 0.96	83.59 $\pm$ 0.71	52.22 $\pm$ 0.87	81.33 $\pm$ 3.14	43.22 $\pm$ 0.21	77.05 $\pm$ 3.83	30.71 $\pm$ 4.56	88.26 $\pm$ 0.60	88.26 $\pm$ 0.60
ceftriaxone	74.81 $\pm$ 3.52	77.04 $\pm$ 3.62	81.13 $\pm$ 1.98	55.63 $\pm$ 4.02	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	87.57 $\pm$ 3.88	68.11 $\pm$ 0.96	83.59 $\pm$ 0.71	52.22 $\pm$ 0.87	81.33 $\pm$ 3.14	43.22 $\pm$ 0.21	77.05 $\pm$ 3.83	30.71 $\pm$ 4.56	88.26 $\pm$ 0.60	88.26 $\pm$ 0.60
ceftriaxone	74.81 $\pm$ 3.52	77.04 $\pm$ 3.62	81.13 $\pm$ 1.98	55.63 $\pm$ 4.02	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	87.57 $\pm$ 3.88	68.11 $\pm$ 0.96	83.59 $\pm$ 0.71	52.22 $\pm$ 0.87	81.33 $\pm$ 3.14	43.22 $\pm$ 0.21	77.05 $\pm$ 3.83	30.71 $\pm$ 4.56	88.26 $\pm$ 0.60	88.26 $\pm$ 0.60
ceftriaxone	74.81 $\pm$ 3.52	77.04 $\pm$ 3.62	81.13 $\pm$ 1.98	55.63 $\pm$ 4.02	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	87.57 $\pm$ 3.88	68.11 $\pm$ 0.96	83.59 $\pm$ 0.71	52.22 $\pm$ 0.87	81.33 $\pm$ 3.14	43.22 $\pm$ 0.21	77.05 $\pm$ 3.83	30.71 $\pm$ 4.56	88.26 $\pm$ 0.60	88.26 $\pm$ 0.60
ceftriaxone	74.81 $\pm$ 3.52	77.04 $\pm$ 3.62	81.13 $\pm$ 1.98	55.63 $\pm$ 4.02	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	87.57 $\pm$ 3.88	68.11 $\pm$ 0.96	83.59 $\pm$ 0.71	52.22 $\pm$ 0.87	81.33 $\pm$ 3.14	43.22 $\pm$ 0.21	77.05 $\pm$ 3.83	30.71 $\pm$ 4.56	88.26 $\pm$ 0.60	88.26 $\pm$ 0.60
ceftriaxone	74.81 $\pm$ 3.52	77.04 $\pm$ 3.62	81.13 $\pm$ 1.98	55.63 $\pm$ 4.02	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	87.57 $\pm$ 3.88	68.11 $\pm$ 0.96	83.59 $\pm$ 0.71	52.22 $\pm$ 0.87	81.33 $\pm$ 3.14	43.22 $\pm$ 0.21	77.05 $\pm$ 3.83	30.71 $\pm$ 4.56	88.26 $\pm$ 0.60	88.26 $\pm$ 0.60
ceftriaxone	74.81 $\pm$ 3.52	77.04 $\pm$ 3.62	81.13 $\pm$ 1.98	55.63 $\pm$ 4.02	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	87.57 $\pm$ 3.88	68.11 $\pm$ 0.96	83.59 $\pm$ 0.71	52.22 $\pm$ 0.87	81.33 $\pm$ 3.14	43.22 $\pm$ 0.21	77.05 $\pm$ 3.83	30.71 $\pm$ 4.56	88.26 $\pm$ 0.60	88.26 $\pm$ 0.60
ceftriaxone	74.81 $\pm$ 3.52	77.04 $\pm$ 3.62	81.13 $\pm$ 1.98	55.63 $\pm$ 4.02	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	87.57 $\pm$ 3.88	68.11 $\pm$ 0.96	83.59 $\pm$ 0.71	52.22 $\pm$ 0.87	81.33 $\pm$ 3.14	43.22 $\pm$ 0.21	77.05 $\pm$ 3.83	30.71 $\pm$ 4.56	88.26 $\pm$ 0.60	88.26 $\pm$ 0.60
ceftriaxone	74.81 $\pm$ 3.52	77.04 $\pm$ 3.62	81.13 $\pm$ 1.98	55.63 $\pm$ 4.02	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	87.57 $\pm$ 3.88	68.11 $\pm$ 0.96	83.59 $\pm$ 0.71	52.22 $\pm$ 0.87	81.33 $\pm$ 3.14	43.22 $\pm$ 0.21	77.05 $\pm$ 3.83	30.71 $\pm$ 4.56	88.26 $\pm$ 0.60	88.26 $\pm$ 0.60
ceftriaxone	74.81 $\pm$ 3.52	77.04 $\pm$ 3.62	81.13 $\pm$ 1.98	55.63 $\pm$ 4.02	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	87.57 $\pm$ 3.88	68.11 $\pm$ 0.96	83.59 $\pm$ 0.71	52.22 $\pm$ 0.87	81.33 $\pm$ 3.14	43.22 $\pm$ 0.21	77.05 $\pm$ 3.83	30.71 $\pm$ 4.56	88.26 $\pm$ 0.60	88.26 $\pm$ 0.60
ceftriaxone	74.81 $\pm$ 3.52	77.04 $\pm$ 3.62	81.13 $\pm$ 1.98	55.63 $\pm$ 4.02	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	86.63 $\pm$ 1.62	81.64 $\pm$ 1.64	87.57 $\pm$ 3.88	68.11 $\pm$ 0.96	83.59 $\pm$ 0.71	52.22 $\pm$ 0.87	81.33 $\pm$ 3.14	43.22 $\pm$ 0.21	77.05 $\pm$ 3.83	30.71 $\pm$ 4.56	88.26 $\pm$ 0.60	88.26 $\pm$ 0.60
ceftriaxone	74.81 $\pm$ 3.52	77.04 $\pm$ 3.62	81.13 $\pm$ 1.98	55.63 $\pm$ 4.02	86.63 $\pm$ 1.62	81.64 $\pm$ 1.6												

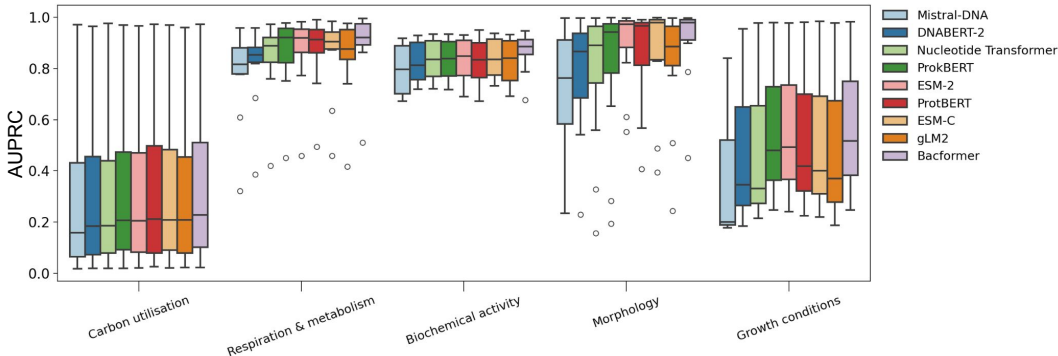
Supplementary Table 14: Per-antibiotic MIC regression prediction results across 3 seeds. Values are mean $\pm$ std. Bold indicates the best mean for each metric.

Antibiotic	Miscral-DNA	DNABERT-2	Nucleotide Transformer	ProkBERT	ESM-2	ESM-4	ProkBERT	gLM2	Bioformer
	Pearson	R ²	Pearson	R ²	Pearson	R ²	Pearson	R ²	Pearson
amoxicillin	0.738 $\pm$ 0.091	45.08 $\pm$ 1.48	0.637 $\pm$ 1.03	46.63 $\pm$ 1.47	0.678 $\pm$ 0.724	48.66 $\pm$ 1.71	0.713 $\pm$ 0.99	50.50 $\pm$ 1.47	71.32 $\pm$ 1.59
amoxicillin-clavulanic acid	46.27 $\pm$ 1.35	21.06 $\pm$ 1.03	48.59 $\pm$ 2.53	23.28 $\pm$ 2.06	52.08 $\pm$ 1.26	26.89 $\pm$ 1.04	51.14 $\pm$ 1.31	57.28 $\pm$ 0.67	50.46 $\pm$ 2.08
ampicillin	46.43 $\pm$ 0.78	20.94 $\pm$ 0.81	48.59 $\pm$ 2.53	23.28 $\pm$ 2.06	52.08 $\pm$ 1.26	26.89 $\pm$ 1.04	51.14 $\pm$ 1.31	57.28 $\pm$ 0.67	32.09 $\pm$ 0.76
ampicillin-sulbactam	40.77 $\pm$ 3.28	16.57 $\pm$ 2.74	41.23 $\pm$ 3.29	16.77 $\pm$ 2.37	45.36 $\pm$ 5.54	20.43 $\pm$ 0.64	42.17 $\pm$ 0.81	47.78 $\pm$ 0.62	31.79 $\pm$ 1.36
azithromycin	79.15 $\pm$ 0.38	62.59 $\pm$ 0.66	79.57 $\pm$ 0.06	63.18 $\pm$ 0.10	80.01 $\pm$ 0.16	63.94 $\pm$ 0.21	80.43 $\pm$ 0.58	80.27 $\pm$ 0.33	22.09 $\pm$ 7.02
aztreonam	72.00 $\pm$ 1.08	50.51 $\pm$ 2.22	72.36 $\pm$ 1.57	51.28 $\pm$ 2.83	74.70 $\pm$ 2.30	55.04 $\pm$ 3.92	78.07 $\pm$ 2.12	79.73 $\pm$ 0.14	64.33 $\pm$ 0.70
cefazolin	46.40 $\pm$ 2.99	23.05 $\pm$ 2.22	49.05 $\pm$ 2.10	24.28 $\pm$ 1.43	53.21 $\pm$ 2.92	28.21 $\pm$ 3.28	53.20 $\pm$ 1.58	52.34 $\pm$ 0.59	36.43 $\pm$ 3.12
cefepime	15.28 $\pm$ 5.95	1.04 $\pm$ 0.84	33.60 $\pm$ 12.14	10.01 $\pm$ 6.23	53.74 $\pm$ 9.12	25.09 $\pm$ 9.05	53.74 $\pm$ 9.12	53.74 $\pm$ 9.12	36.34 $\pm$ 2.49
ceftriaxone	9.42 $\pm$ 6.07	-2.14 $\pm$ 5.56	16.40 $\pm$ 8.74	0.97 $\pm$ 4.55	41.05 $\pm$ 2.51	25.09 $\pm$ 9.05	53.74 $\pm$ 9.12	53.74 $\pm$ 9.12	36.34 $\pm$ 2.49
cefuroxime	54.85 $\pm$ 1.67	29.82 $\pm$ 1.66	60.37 $\pm$ 0.97	36.04 $\pm$ 0.86	63.07 $\pm$ 0.88	39.67 $\pm$ 1.15	62.02 $\pm$ 1.02	61.36 $\pm$ 0.39	52.60 $\pm$ 1.88
cefuroxime sodium	32.78 $\pm$ 4.16	10.27 $\pm$ 2.16	32.72 $\pm$ 3.35	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56 $\pm$ 3.02	-0.02 $\pm$ 3.00	10.00 $\pm$ 1.46	10.00 $\pm$ 1.46	40.53 $\pm$ 4.54	15.11 $\pm$ 2.73	42.37 $\pm$ 0.25	38.75 $\pm$ 0.69	66.00 $\pm$ 1.27
cefuroxime sodium	4.56								

C.6 PHENOTYPIC TRAITS PREDICTION

We measure the AUPRC across phenotype groups (Supp. Fig. 13) as well as overall performance across metrics and all phenotypic traits (Supp. Table 15), and groups (Supp. Table 16). Finally, we include the results for each individual phenotype (Supp. Table 17 & 18). Analyzing performance across phenotype groups and metrics (Supp. Fig. 13 & Supp. Table 16), bLM-Bacformer achieves the best or combined best result across all groups and metrics, showing the benefits of incorporating genomic interactions into phenotype prediction models. Across metrics and phenotypes, Bacformer achieves the best result on 80 phenotypes on AUROC, 78 on AUPRC and 86 on F1 out of a total of 139 phenotypes. Therefore, showing that there is a considerable variation between phenotypes and one should choose a model specific for a phenotype. We believe this is due to the variable number of labels available for a phenotype as well as the inherent differences in the phenotypes themselves, which make some phenotypes easier to predict using a computational approach than others.

Predicting phenotypes is often a very challenging task which includes understanding the effect of mutations and multi-level genomic interactions. However, accurately predicting phenotypes could allow us to engineer genomes for a desired purpose, such as sustainable bioproduction, thus having potentially massive positive impact. We believe that next-generation models should consider the genomic context and incorporate the prior knowledge, such as genome-scale metabolic models to make the most out of available data for a given phenotype.



Supplementary Figure 13: AUPRC across diverse phenotypic traits groups and methods. The box spans the inter-quartile range with a line marking the median value. The results were macro-averaged across classes for each phenotype.

Supplementary Table 15: Overall phenotypic traits prediction performance across all phenotypes. Values are mean $\pm$ standard deviation across 5 seeds. The results were macro-averaged across classes for each phenotype.

Method	AUROC	AUPRC	F1
Mistral-DNA	60.34 $\pm$ 9.81	39.72 $\pm$ 33.09	47.25 $\pm$ 6.67
DNABERT-2	63.83 $\pm$ 11.06	42.10 $\pm$ 34.24	49.42 $\pm$ 8.99
Nucleotide Transformer	65.98 $\pm$ 11.45	43.22 $\pm$ 34.85	52.04 $\pm$ 11.07
ProkBERT	67.64 $\pm$ 11.87	44.80 $\pm$ 34.72	52.03 $\pm$ 11.56
ESM-2	68.71 $\pm$ 12.36	45.61 $\pm$ 35.45	53.10 $\pm$ 11.90
ESM-C	69.07 $\pm$ 11.98	45.54 $\pm$ 35.38	52.28 $\pm$ 11.39
ProtBERT	68.47 $\pm$ 11.97	45.47 $\pm$ 34.89	53.71 $\pm$ 11.65
gLM2	65.82 $\pm$ 11.41	44.00 $\pm$ 34.40	51.11 $\pm$ 10.86
Bacformer	71.34 $\pm$ 12.91	47.80 $\pm$ 35.99	56.14 $\pm$ 13.19

Supplementary Table 16: Phenotype-group prediction results; *mean $\pm$ standard deviation* over five random seeds. The highest score in each column is shown in bold. Results are macro-averaged across classes for each phenotype.

Method	Biochemical activity			Carbon utilisation			Growth conditions			Morphology			Respiration & metabolism		
	AUROC	AUPRC	F1	AUROC	AUPRC	F1	AUROC	AUPRC	F1	AUROC	AUPRC	F1	AUROC	AUPRC	F1
Mistral-DNA	54.67 $\pm$ 12.73	79.60 $\pm$ 10.06	44.48 $\pm$ 2.40	58.17 $\pm$ 7.25	27.95 $\pm$ 27.92	46.67 $\pm$ 3.25	63.73 $\pm$ 4.34	40.58 $\pm$ 37.70	36.80 $\pm$ 17.74	70.07 $\pm$ 11.96	72.05 $\pm$ 23.63	52.81 $\pm$ 9.92	71.79 $\pm$ 11.54	77.30 $\pm$ 18.47	50.88 $\pm$ 14.54
DNAERT-2	59.88 $\pm$ 9.16	82.39 $\pm$ 8.55	46.19 $\pm$ 6.05	60.46 $\pm$ 7.60	29.24 $\pm$ 28.29	47.63 $\pm$ 3.90	71.22 $\pm$ 17.07	49.41 $\pm$ 40.69	45.89 $\pm$ 27.95	79.15 $\pm$ 12.39	79.75 $\pm$ 22.57	56.83 $\pm$ 14.23	76.88 $\pm$ 11.06	80.97 $\pm$ 16.72	60.52 $\pm$ 15.97
Nucleotide Transformer	63.74 $\pm$ 8.59	83.40 $\pm$ 8.47	50.62 $\pm$ 4.88	62.07 $\pm$ 7.58	30.06 $\pm$ 28.72	49.20 $\pm$ 4.03	79.33 $\pm$ 13.37	50.82 $\pm$ 41.12	51.73 $\pm$ 33.88	79.63 $\pm$ 14.10	78.94 $\pm$ 26.27	61.55 $\pm$ 20.66	81.74 $\pm$ 10.26	83.90 $\pm$ 16.08	66.91 $\pm$ 14.99
ProBERT	61.84 $\pm$ 8.89	83.41 $\pm$ 8.29	49.52 $\pm$ 5.28	63.50 $\pm$ 7.72	31.50 $\pm$ 28.50	49.34 $\pm$ 4.49	86.10 $\pm$ 8.33	56.82 $\pm$ 37.42	49.07 $\pm$ 29.41	82.76 $\pm$ 12.95	81.29 $\pm$ 26.40	61.87 $\pm$ 23.45	84.91 $\pm$ 8.16	85.55 $\pm$ 16.14	66.74 $\pm$ 14.85
ESM-2	62.86 $\pm$ 12.75	83.30 $\pm$ 9.29	51.30 $\pm$ 9.54	64.60 $\pm$ 8.18	31.77 $\pm$ 29.34	49.38 $\pm$ 4.58	81.70 $\pm$ 11.20	56.97 $\pm$ 37.48	52.70 $\pm$ 29.63	87.97 $\pm$ 10.47	89.36 $\pm$ 14.86	71.41 $\pm$ 18.24	83.46 $\pm$ 11.48	86.63 $\pm$ 15.62	66.11 $\pm$ 16.31
ESM-C	62.71 $\pm$ 11.90	83.90 $\pm$ 8.16	50.82 $\pm$ 9.24	65.37 $\pm$ 8.16	32.20 $\pm$ 29.54	49.27 $\pm$ 4.24	78.71 $\pm$ 15.27	53.44 $\pm$ 40.05	55.72 $\pm$ 35.81	88.05 $\pm$ 9.67	87.26 $\pm$ 20.11	67.68 $\pm$ 20.43	81.79 $\pm$ 12.13	85.05 $\pm$ 16.95	61.05 $\pm$ 14.90
ProBERT	61.93 $\pm$ 13.25	82.30 $\pm$ 10.18	50.78 $\pm$ 7.31	64.86 $\pm$ 8.29	32.22 $\pm$ 29.14	50.30 $\pm$ 5.05	82.07 $\pm$ 11.45	54.06 $\pm$ 39.32	54.87 $\pm$ 33.97	83.27 $\pm$ 13.36	86.10 $\pm$ 18.70	68.60 $\pm$ 18.70	84.71 $\pm$ 8.40	86.48 $\pm$ 14.80	68.69 $\pm$ 14.17
gLM2	60.39 $\pm$ 8.39	82.71 $\pm$ 9.45	47.76 $\pm$ 4.82	62.78 $\pm$ 7.22	30.97 $\pm$ 28.27	48.65 $\pm$ 4.06	71.55 $\pm$ 26.57	51.10 $\pm$ 41.36	45.49 $\pm$ 27.39	77.44 $\pm$ 16.90	82.79 $\pm$ 21.88	62.18 $\pm$ 20.77	82.26 $\pm$ 9.24	84.22 $\pm$ 16.68	64.43 $\pm$ 16.91
Bacformer	67.42 $\pm$ 16.45	86.32 $\pm$ 9.83	55.09 $\pm$ 11.26	66.94 $\pm$ 8.94	33.94 $\pm$ 30.24	51.99 $\pm$ 6.44	86.18 $\pm$ 9.85	58.25 $\pm$ 37.22	59.53 $\pm$ 33.08	89.53 $\pm$ 10.01	91.13 $\pm$ 15.07	70.20 $\pm$ 19.26	88.18 $\pm$ 9.31	89.19 $\pm$ 14.12	77.15 $\pm$ 15.49

Supplementary Table 17: Per-phenotype performance (1/2). Values are mean $\pm$ standard deviation across 5 random seeds. Bold highlights the best mean per metric. The results were macro-averaged across classes for each phenotype.

[illegible]

Supplementary Table 18: Per-phenotype performance (2/2). Values are mean $\pm$ standard deviation across 5 random seeds. Bold highlights the best mean per metric. The results were macro-averaged across classes for each phenotype.

[illegible]

D BROADER IMPACT & LIMITATIONS

Broader impact BacBench provides the first public, multi-task testbed that spans gene-, system- and genome-scale prediction problems over 67k genomes from 17.6k species. By standardising data splits, evaluation code and baseline embeddings, it lowers the entry barrier for machine-learning researchers who lack domain-specific pipelines yet want to work on microbial genomics. In the near term, more reliable essential-gene or antibiotic-resistance predictors could shorten drug-development cycles and inform stewardship policies, while better phenotype-from-genome models will accelerate the search for chassis strains that sequester carbon, degrade waste or synthesise valuable biochemicals. Because the benchmark emphasises cross-species generalization, methods that succeed on BacBench are naturally suited to poorly studied or newly sequenced taxa, helping global health laboratories track emerging pathogens even when only draft assemblies are available. Finally, releasing all data under permissive licences and exposing a HuggingFace hub invites continual community contributions, which should foster an open, comparative culture similar to computer-vision or NLP benchmarks and, in turn, drive rapid, reproducible advances in microbial bio-AI research.

Limitations Despite its breadth, BacBench still samples an uneven slice of bacterial diversity: phenotypic-trait labels cluster heavily around medically important genera, and some antibiotic classes remain sparsely represented, which could bias models towards well-studied lineages and mechanisms. Tasks that matter for ecology and biotechnology—horizontal-gene-transfer detection, host–phage interaction, metabolic-flux prediction or transcriptome conditioning—are absent, so performance on BacBench should not be interpreted as general mastery of bacterial genomics. Moreover, the benchmark inherits experimental noise from upstream databases: STRING DB interaction scores mix heterogeneous evidence; operon annotations are incomplete; and phenotype labels amalgamate disparate growth protocols, introducing label uncertainty that caps achievable accuracy. Finally, computing embeddings for every update is resource-intensive, which may hinder participation from groups without access to multi-GPU servers, although smaller surrogate splits are planned for future releases.