



A sequence-based classifier distinguishes phenotype-associated genes from other gene models in plants

Nikee Shrestha, Zhongjie Ji, Xiuru Dai, et al.

Genome Res. published online July 14, 2026

Access the most recent version at doi:[10.1101/gr.281802.125](https://doi.org/10.1101/gr.281802.125)

P<P	Published online July 14, 2026 in advance of the print journal.
Accepted Manuscript	Peer-reviewed and accepted for publication but not copyedited or typeset; accepted manuscript is likely to differ from the final, published version.
Open Access	Freely available online through the <i>Genome Research</i> Open Access option.
Creative Commons License	This manuscript is Open Access. This article, published in <i>Genome Research</i> , is available under a Creative Commons License (Attribution 4.0 International license), as described at http://creativecommons.org/licenses/by/4.0/ .
Email Alerting Service	Receive free email alerts when new articles cite this article - sign up in the box at the top right corner of the article or click here .

Advance online articles have been peer reviewed and accepted for publication but have not yet appeared in the paper journal (edited, typeset versions may be posted when available prior to final publication). Advance online articles are citable and establish publication priority; they are indexed by PubMed from initial publication. Citations to Advance online articles must include the digital object identifier (DOIs) and date of initial publication.

To subscribe to *Genome Research* go to:
<https://genome.cshlp.org/subscriptions>

Published by Cold Spring Harbor Laboratory Press

A sequence-based classifier distinguishes phenotype-associated genes from other gene models in plants

Nikee Shrestha^{1,2}, Zhongjie Ji¹, Xiuru Dai³, Pinghua Li³, and James C. Schnable^{*,1,2}

¹Center for Plant Science Innovation, University of Nebraska-Lincoln, Lincoln, NE, 68503, USA

²Department of Agronomy and Horticulture, University of Nebraska-Lincoln, Lincoln, NE, 68583, USA

³State Key Laboratory of Crop Biology, College of Agronomic Sciences, Shandong Agricultural University, Tai'an, Shandong, 271018, China

*schnable@unl.edu

ABSTRACT

Only a small fraction of annotated plant genes possess experimentally validated associations with specific phenotypes. Phenotype associated genes have distinct structural, molecular, and evolutionary characteristics compared to non-validated gene models. Here, we develop a simple classifier that uses sequence and evolutionary features which can be generated for any species with an annotated reference genome assembly, to accurately distinguish phenotype associated genes from both the overall population of annotated gene models and a specific set of genes identified as being tolerant of premature stop mutations. A model trained solely on genes from maize (*Zea mays*) identify and prioritize rice (*Oryza sativa*) and *Arabidopsis* (*Arabidopsis thaliana*) genes that are highly enriched in genes with experimentally validated links to phenotypes in both of these evolutionarily distant species. Gene models predicted to have a higher probability of being linked to phenotypes display patterns consistent with known biological properties of phenotype associated genes. Notably, the sets of genes predicted to have

1 a high probability of being linked to phenotype variation do not consist exclusively of well-
2 characterized gene families but included many uncharacterized gene families carrying domains
3 of unknown function. The quantitative scores generated by this model offer a valuable resource
4 for prioritizing and exploring the vast number of uncharacterized gene models in plants, reducing
5 the risk of failure in future reverse genetics efforts and potentially accelerating gene discovery
6 and functional annotation in crops.

7 INTRODUCTION

8 The genomes of typical multicellular eukaryotes contain tens of thousands of gene model
9 annotations. However, few of these genes have been the subject of specific genetic or molecular
10 investigations. Indeed, in *Arabidopsis* less than 10% of genes have been linked to specific
11 phenotypes (Lloyd and Meinke, 2012); likewise, in rice approximately 10% of gene models are
12 associated with phenotypes (Huang et al., 2022), and in maize this is less than 1% (Schnable and
13 Freeling, 2011). The situation in other plants is even worse (Oellrich et al., 2015). This is not
14 unique to plants; even in the human genome many annotated genes have not been the subject of
15 detailed investigation by even a single study (Sinha et al., 2000), including hundreds of genes
16 known to be essential for cell survival (Wang et al., 2015). Research continues to focus on the
17 same well-studied genes (Kustatscher et al., 2022a,b), thereby hindering our understanding of the
18 core functionalities of eukaryotes and the mechanisms responsible for lineage-specific processes
19 and phenotypes.

20 The process of functionally characterizing less-studied gene models or generating new
21 hypotheses about gene-phenotype associations is time-intensive and high-risk. These efforts can
22 take years and frequently yield inconclusive or null results, creating a very real career risk for
23 researchers. These limitations constrain both reverse genetics approaches and strategies that

leverage natural genetic diversity. Marker trait association signals that fall near already known genes tend to be viewed as confirmatory. In contrast, quantitative trait nucleotides located next to gene models not previously implicated in a phenotype are less pursued, largely because functional validation requires substantial time, resources, and carries a high risk of yielding null results. Similar obstacles arise when interpreting results from differential gene expression analyses, co-expression networks, and other multi-omics approaches. In contrast, once a gene has been shown to influence at least one phenotype, the risk of investigation declines, further amplifying attention toward genes with existing functional evidence while discouraging exploratory work on the vast number of uncharacterized gene models (Kwon et al., 2024). For example, the maize teosinte branched1 (*tb1*) gene plays an important role in reducing tillering in domesticated maize (Doebley et al., 1990, 1997). The association of *tb1* with important domestication phenotypes in maize drove the study of this gene to understand its biological role in maize and in other species (Igartua et al., 2020) with at least 140 additional studies examining *tb1* since the original characterization per MaizeGDB (Woodhouse et al., 2025). In principle characterization of a gene likely also reduces the perceived risk of studying other members of the same gene family, resulting on concentration of investigation not only on specific genes but specific gene families, although the latter effect has not been rigorously evaluated in plants. Being able to predict, in advance, which genes are likely to influence phenotypes would reduce the risk of investigating currently uncharacterized or under characterized genes.

Previous studies have shown that classifiers can distinguish genes associated with different classes of phenotypes. For example, prior work has separated lethal from nonlethal genes in *Arabidopsis*, distinguished essential from nonessential genes across six model eukaryotes with cross-species applicability (Lloyd et al., 2015, Beder et al., 2021), and classified

Arabidopsis genes involved in primary versus specialized metabolism (Moore et al., 2019) or even prediction of gene-to-phenotype associations to prioritize candidate genes for traits (Lee et al., 2010). Features used in these studies typically include both sequence-derived properties and diverse omics data collected across tissues and experiments. Beyond distinguishing biologically meaningful gene sets, previous work has also shown that annotation errors can be identified computationally, including in maize using evolutionary conservation and in *Pisum sativum* using nucleotide and amino acid composition features (Upadhyay et al., 2022; Schulz et al., 2025). It has been shown that genes already associated with phenotypes differ from other gene models across a broad range of characteristics, and these differences appear to be consistent regardless of whether the associated phenotype is lethal, constitutive, conditional, or detectable only at the biochemical or molecular level (Schnable, 2020). This pattern suggests the existence of a distinct class of gene models that are either truly unrelated to organismal phenotype or, at a minimum, unlikely to be linked to phenotype using current experimental approaches. Studying this class of genes is challenging. The set of gene models lacking a known phenotype comprises a mixture of both genes, which could be linked to phenotype if studied, and those truly unrelated to determining the observable phenotypes of an organism. The lack of incentives to publish null results makes it difficult to specifically identify the subset of gene models where researchers have looked for links to phenotypes and found none. In this study, we sought to evaluate how well genes known to be associated with studied phenotypes could be distinguished from the background set of annotated gene models, whether it was possible to distinguish them based solely on features which could be derived from genomic sequence, and whether models trained in a single genetic model species are likely to perform well in other plant species.

Results

We identified 272 maize gene models linked to phenotypic outcomes based on evidence from loss-of-function alleles and named this set as phenotype-associated genes (PAGs) (Supplemental_Dataset_S1). Genes where loss of function alleles have been studied and were not linked to observed phenotypic consequences were poorly documented and it was not possible to identify a large enough set to employ as a negative set for model training. As an alternative approach to identifying genes in which loss-of-function alleles were tolerated, we identified a set of gene models where alleles creating premature stop codons were widespread using whole-genome resequencing data from 1,515 maize genotypes that include open pollinated varieties and wild relatives. This resulted in a second set of 713 gene models with variants causing a premature stop codon, i.e., premature stop enriched genes (PSEs). There was no overlap between the set of PAGs and the set of PSEs ($E = 5$ gene models, $p\text{-value} = 0.017$; Fisher's exact test). In many cases, the alleles that produced premature stop codons were relatively rare, defined as premature stop enriched-rare (PSE-R) genes, but in 162 cases, the frequency of the allele that resulted in a premature stop codon was greater than 0.25 that excludes across the re-sequenced population used to discover these variants (Figure 1A). We focused on these 162 genes that carry common premature stop codon variants in downstream analyses, as rarer premature stop codons may be more likely to represent alleles which would impact fitness when homozygous but are still represented in our population datasets via data from open pollinated varieties or wild relatives. These gene models were supported by transcriptional data and, frequently, by proteomic and comparative genomic evidence, and thus do not simply represent artifacts produced by ab initio gene annotation software. Below, we refer to this set of gene models as premature stop enriched-common (PSE-C) genes. Both PAGs and PSE-Cs were depleted in the category of gene models lacking evidence of expression in any tissue. Of the 1,215 gene models

1 meeting these criteria, 3 were PAGs (8 expected), and zero were PSE-Cs (5 expected). We
2 compared PAGs and PSE-Cs with common premature stop codon alleles across 2,604 features
3 including amino acid and nucleotide composition (as well as dinucleotide frequency, etc),
4 structural characterized of gene models (e.g. sequence length, number of exons), predicted
5 protein functional characteristics (e.g. hydrophobicity), gene expression and protein abundance
6 data, and features derived from chromatin, co-expression, sequence variation, or transcription
7 factor binding datasets (See Methods, Supplemental_Dataset_S2). Both PAGs and PSE-Cs were
8 less common among non-expressed genes than expressed genes (3 non-expressed PAGs, 8
9 expected; 0 non-expressed PSE-Cs, 5 expected). PAGs and PSE-Cs exhibited statistically
10 significant differences for about 44% of the time (p -value < 0.05 , two-tailed Mann-Whitney U
11 test, Supplemental_Dataset_S2), with 16% of these differences exceeded a more stringent
12 Bonferroni corrected p -value (p -value $< 0.05/2,604$ tested features). PAGs and PSE-Cs exhibited
13 significant differences in gene length, average protein abundance, the number of tissues
14 expressing a gene model (Supplemental Figure S1), and syntenic conservation in sorghum (p -
15 value < 0.001 , $2.2 * 10^{-16}$; Fisher's exact test). These results are consistent with previously
16 reported differences between genes associated with mutant phenotypes and the overall set of
17 annotated gene models (Schnable, 2020) and also between lethal and non-lethal phenotype
18 associated *Arabidopsis* genes (Lloyd et al., 2015). No significant difference was observed
19 between the number of exons in PAGs and PSE-Cs, another trait previously reported to differ
20 between genes associated with mutant phenotypes and other annotated gene models (p -
21 value=0.53, two-tailed Mann-Whitney U test, Figure 1C). PAGs had significantly longer coding
22 sequences than PSE-Cs, inconsistent with the possibility that the enrichment of premature stop
23 codons resulting from a larger number of potential sites where a mutation could introduce a stop

1 codon. Many other features where PAGs and PSE-Cs significantly differed had not previously
 2 been linked to the likelihood of being associated with organismal phenotypes. These include the
 3 length of the 3' untranslated region (UTR) (PAGs > PSE-Cs p-value < 0.001 [8.13×10^{-15}], two-
 4 tailed Mann-Whitney *U* test, Cohen's $d=0.64$, Figure 1D), number of isoforms (PAGs > PSE-Cs
 5 p-value= 0.003, two-tailed Mann-Whitney *U* test, Cohen's $d=0.23$, Figure 1E), and GC content
 6 in the nucleotide sequence of the gene model (PAGs > PSE-Cs p-value < 0.001 [6.70×10^{-5}],
 7 two-tailed Mann-Whitney *U* test, Cohen's $d=0.53$, Figure 1F). Maize genes exhibit a bimodal
 8 distribution of GC content (Alexandrov et al., 2009). The distributions for both PAGs and PSE-
 9 C genes both reflected the same bimodal distribution observed for maize genes as a whole,
 10 although PSE-Cs included a higher ratio of the lower GC content mode relative to n PAGs
 11 (Supplemental Figure S2). PAGs appeared to be more distinct from the overall population of
 12 annotated gene models than PSE-Cs. Significant differences between PAGs and non-PAG, non-
 13 PSE-C genes were detected for 54% of quantitative features (mean Cohen's $d = 0.05$; median
 14 Cohen's $d = 0.01$), with 19% remaining significant after Bonferroni correction (p-value <
 15 $0.05/2,604$ tested features). By contrast, significant differences between PSE-Cs and non-PAG,
 16 non-PSE-C genes were detected for only 28% of quantitative features (mean Cohen's $d = 0.008$;
 17 median Cohen's $d = 0.002$), with only 3% remaining significant after Bonferroni correction (p-
 18 value < $0.05/2,604$ tested features) (Supplemental_Dataset_S2). Collectively, these results
 19 demonstrate that PSE-Cs are more similar to the rest of the gene models than PAGs, and that
 20 PAGs are significantly different from both PSE-Cs and the rest of the average gene models for
 21 both previously tested and untested features.

22 Despite the significant differences between PAGs and PSE-Cs for many individual
 23 features, these two populations of genes exhibited only modest separation on the first principal

1 component calculated from the full set of gene features (Figure 1B). This indicates that there
2 may be multiple distinct patterns of genes whose functions are likely (or unlikely) to be
3 associated with phenotypes. This result also contrasts with the previously reported clear
4 separation between phenotype-associated genes and other annotated gene models when using a
5 smaller, curated feature set (Schnable, 2020). We trained a random forest model, which tends to
6 be more robust to nonlinear relationships, to distinguish PAGs and PSE-Cs using all 2,675 gene
7 features (2,604 continuous or numeric features + 71 binary features). Model performance was
8 evaluated using receiver operating characteristic (ROC) curves and the area under the curve
9 (AUC) using a holdout dataset of 48 gene models (PAGs=29, PSE-Cs=19). ROC curves visually
10 summarize the trade-off between true and false positive rates across different thresholds, while
11 the ROC-AUC provides a single numeric metric which quantifies performance across all
12 thresholds. The model exhibited strong predictive performance achieving an ROC-AUC of 0.87
13 on the hold-out test gene sets (Figure 2A). The prediction probabilities assigned to PAGs
14 (mean=0.89) and PSE-Cs (mean=0.17) were each significantly different from those of
15 uncharacterized annotated gene models (mean=0.50, p -value<0.001 for both comparisons, two-
16 tailed Mann-Whitney U test, Supplemental_Dataset_S3). One third of uncharacterized maize
17 gene models were assigned lower scores than any of the genes in our dataset already linked to
18 phenotypes, implying that a substantial fraction of annotated gene models representing sequences
19 unlikely to produce publishable or easy to study phenotypes (Figure2B). Above this threshold,
20 many additional annotated genes were assigned scores lower than the highest scoring PSE-C.
21 Both of these results are consistent with a substantial fraction of annotated gene models being
22 unlikely to produce publishable or easy to study phenotypes when disrupted. The ability of a

1 random forest model to achieve strong predictive performance on hold-out test sets of PSE-Cs
2 and PAGs supports the substantial differentiation of these populations of gene models.

3 We next sought to determine how many, and which types of features were necessary to
4 distinguish these populations of genes. The set of features describing maize gene models can be
5 separated into genomic-sequence-derived and non-genomic sequence derived features. Genomic-
6 sequence-derived features can be directly extracted or calculated, given only an assembled
7 genome sequence and a structural gene annotation file. By contrast, non-genomic sequence
8 derived features, including gene expression data, protein abundance, chromatin accessibility
9 profiling, and assays of transcription factor binding, require additional experiments and are
10 typically challenging to generate similarly across experiments conducted in multiple species.

11 Among the 100 features with the highest Gini importance in the trained PAG/PSE-C
12 classification model, 12 were derived from non genomic sequence or expression data. Gini
13 importance quantifies how much each feature contributes to reducing classification error across
14 the decision trees in the model. Only one of the top 10 most important features was not derived
15 from genomic sequence, namely the maximum Fragments Per Kilobase of transcript per Million
16 mapped reads (FPKM) observed for a gene across all tissue samples (Supplemental_Dataset_S4).
17 Syntenic conservation of a gene model in sorghum appeared to be the most important feature
18 (Supplemental_Dataset_S4). The top 100 features also included gene structural attributes,
19 particularly 5' and 3' untranslated region lengths, as well as numerous protein descriptors
20 reflecting amino acid composition, residue ordering, physicochemical properties, and local
21 sequence architecture. Other prominent features included *k*-mer frequencies, frequency of
22 occurrence of each amino acids (codon usage), the number of times a codon appears in a gene
23 divided by the number of expected occurrences under equal codon usage (relative synonymous

codon usage), and GC content at the third codon position consistent with differences between PAGs and PSE-Cs (Figure 1F). A random forest model trained solely on genomic and comparative genomic derived features achieved comparable performance, with an ROC curve area of 0.86 (Figure 2A, B). Model performance remained consistent with an ROC curve area of 0.87 when the feature set was restricted to only 100 genomic sequence derived features with the highest Gini importance (Figure 2A, B). Consistent with previous models, synteny remained the most important feature. The probability of a gene in the top 10% (threshold=0.88, n=3,904) of predicted scores generated by the top hundred-feature genomic sequence derived data model being linked to a phenotype was estimated to be approximately 18.3 times higher than that of a gene with a score in the bottom 50% (threshold=0.88, n=19,516) based on evaluation of the hold-out test data. These results suggest that it may be feasible to build high-performing models based solely on feature sets that can be generated consistently across different species.

The model trained solely on genome sequence and comparative genome derived features was able to separate the prediction probability distribution of premature stop enriched-common (PSE-C) maize gene models from phenotype associated genes in maize, with uncharacterized gene models having overlapping distribution with both extreme classes (Figure 3A). We next tested how well our model translates to other plant species. The genetic characterization of several plant species, notably rice and *Arabidopsis*, also includes information on the phenotypic consequences of disrupting comparable or greater numbers of genes compared to the maize dataset used in this study. However, in most plant species, with sequenced genomes, few or no genes have been directly linked to mutant phenotypes. We assessed the transferability of a model trained on maize data to other species using a model trained on 99 of the top 100 maize features, excluding one feature that could not be easily generated for rice gene models (see Methods), to

1 predict the probability that annotated rice genes would be associated with phenotypic
2 consequences. From the set of approximately 4,500 cloned rice genes in the funRiceGenes
3 database (Huang et al., 2022), we manually curated a non-exhaustive set of 100 which met the
4 same phenotype criteria applied to maize. The 96 rice PAGs which could be mapped to MSU v7
5 gene models included 72 which were conserved at syntenic locations in maize and 24 which
6 lacked maize syntenic orthologs. However, in only seven cases was the maize syntenic ortholog
7 of a rice PAG also classified as a PAG in maize. The 96 rice PAGs were assigned significantly
8 higher probabilities of being associated with a phenotype than with the background set of
9 annotated gene models in the rice genome (Figure 3B, $p\text{-value} < 0.001 [5.27 * 10^{-19}]$, two-tailed
10 Mann-Whitney U test, Supplemental_Dataset_S5). The probability of a rice gene in the top 10%
11 (threshold=0.63, $n=5,564$) of scores being linked to a phenotype was eleven times higher than
12 that of a gene in the bottom 50% (threshold=0.45, $n=27,015$). The divergence of the lineages
13 leading to maize and rice is estimated to have occurred approximately 50 million years ago,
14 roughly one-third the time since the estimated age of the most recent common ancestor of maize
15 and *Arabidopsis* (~160 million years ago) (Kumar et al., 2017). The *Arabidopsis* genome is also
16 less than 10% the size of the maize genome, with substantially lower transposon content. A
17 model trained solely on maize data using the same 99 features assigned significantly higher
18 probabilities of being associated with a phenotype to a set of *Arabidopsis* PAGs than to the
19 background set of annotated gene models in the *Arabidopsis* genome (Figure 3C, $p\text{-value} <$
20 $0.001 [1.12 * 10^{-176}]$, two-tailed Mann-Whitney U test, Supplemental_Dataset_S5). The
21 prediction probability was not different between essential, morphological, cellular and
22 biochemical *Arabidopsis* PAGs and conditional *Arabidopsis* PAGs ($p\text{-value}=0.77$, two-tailed
23 Mann-Whitney U test, Supplemental Figure S3). The probability of an *Arabidopsis* gene model

1 in the top 10% (threshold=0.52, $n=2,806$) of scores being linked to a phenotype was four times
2 higher than a gene in the bottom 50% (threshold=0.41, $n=13,227$). These results confirm the
3 notion that a model trained solely on maize data can still achieve both statistically significant and
4 operationally meaningful performance in prioritizing genes for phenotypic characterization when
5 applied to a distantly related angiosperm with a radically different genome architecture.

6 Cross-validation and testing on a randomly selected set of hold-out test data are both
7 widely used approaches to estimate the performance of machine learning models. However, the
8 accuracy of machine learning models can sometimes decline when evaluated on truly out of
9 sample datasets relative to estimates based on held out subsets of the same dataset used to train
10 the initial model (Guan et al., 2025; Lei et al., 2025). We generated a new, curated set of maize
11 PAGs from literature published after our original PAG dataset was finalized (Supplemental Data
12 Set S7). The estimated probabilities of being being linked to phenotype of these new out-of-
13 sample PAGs ($n=27$) were significantly higher than those of a typical uncharacterized gene (p -
14 value=0.002, two-tailed Mann-Whitney U test, Figure 3D). The average estimated probability
15 assigned to the out of sample PAG genes was modestly lower than the probabilities assigned to
16 PAGs from the original hold-out test dataset ($\mu=0.68$ vs $\mu=0.79$). However, the estimated
17 probability that a gene in the top 10% of the prediction scores is associated with a phenotype was
18 still approximately 13 times as high as the probability for a gene in the bottom 50% of scores.
19 This ratio was estimated to be 18.3 times using the original hold-out test genes. These results
20 demonstrate that the prediction model retains strong predictive performance even when applied
21 to truly out-of-sample datasets.

22 The gene models predicted to be linked to phenotypes are supported by other data types
23 associated with purifying selection or the consequences of mutations. Gene models in the bottom

quartile (threshold=0.27, n=9,759) of probability of being associated with a phenotype exhibited approximately 4.4 times higher odds of containing a premature stop codon compared to those in the top quartile (threshold=0.76, n=9,759, p-value < 2.2×10^{-16} ; Fisher's exact test) (Figure 4A, B). Variants introducing premature stop codons in bottom-quartile gene models were observed at higher minor allele frequencies than variants introducing premature stop codons in top-quartile genes (p-value=0.01 two-tailed Mann-Whitney *U* test). Based on alignments to orthologous genes in sorghum, bottom-quartile genes exhibited higher rates of sequence evolution at both the DNA sequence and protein sequence level with elevated rates of synonymous substitutions (median K_s , 0.53 vs 0.28, p-value < 0.001), elevated non-synonymous substitutions (median K_a , 0.19 vs 0.03, p-value < 0.001) and elevated K_a/K_s (median K_a/K_s , 0.14 vs 0.20, p-value < 0.001), consistent with relaxed purifying selection or a mixture of purifying and positive selection. Bottom-quartile gene models showed higher CG and CHG methylation levels across annotated gene regions than top-quartile gene models. In contrast, CHH methylation across annotated gene regions was higher in top-quartile gene models than in bottom-quartile gene models, consistent with the bottom-quartile set not being primarily composed of transposon fragments. Variant frequency per kilobase of coding sequence (CDS) was higher in top-quartile gene models than bottom quartile gene models, but this effect was driven by a higher prevalence of variants predicted to have small (low) effects on protein function, typically synonymous variants, and the proportion of variants predicted to have moderate (p-value < 0.001, two-tailed Mann-Whitney *U* test) or strong (high) effects on protein function (p-value < 0.001, two-tailed Mann-Whitney *U* test) were significantly higher in bottom-quartile gene models (Figure 4C, D). The apparent inconsistency between higher K_a and K_s ratios both being observed in bottom quartile genes and a higher overall variant frequency being observed in top quartile genes may have resulted from

the subset of genes which lacked syntenic orthologs in sorghum. These genes were still included in estimates of variant frequency but it was not possible to estimate substitution rates for these genes. A similar pattern may also arise because bottom-quartile genes were more likely to show presence-absence variation. When a gene is present in some maize lines but absent in others, SNP and InDel calls within that gene tend to have higher levels of missing data and are therefore more likely to be removed during quality filtering. We employed PlantCaduceus, a DNA language model that has learned evolutionary conservation patterns in angiosperm genomes (Zhai et al., 2025) to estimate the effects of DNA substitutions without the need to align against genes in related species. Genetic variants in the gene body or CDS of top-quartile gene models exhibited significantly higher predicted likelihoods of being deleterious (based on average absolute zero-shot scores from PlantCaduceus) than bottom-quartile gene models ($p\text{-value} < 1 * 10^{-200}$, two-tailed Mann-Whitney U test, Figure 4E). By contrast, variants present in UTRs or introns exhibited a small but significant bias in the opposite direction, with variants in these regions of bottom-quartile genes assigned marginally higher absolute zero-shot scores by PlantCaduceus than those in top-quartile genes ($p\text{-value} < 1 * 10^{-20}$, two-tailed Mann-Whitney U test, Figure 4E).

To assess whether the model was prioritizing members of well-annotated gene families, we examined Gene Ontology (GO) term enrichment analysis among genes from top quartile and bottom quartile of the distribution, constraining GO enrichment analysis to the population of only those genes assigned at least one GO term. An average top quartile gene models were assigned 17 GO term per gene model and bottom quartile gene models were assigned 6 GO terms per gene model (minimum GO term assigned to a gene model: 1, maximum GO term assigned to a gene model: 297, Supplementary Figure S4). GO enrichment analysis showed that

top-quartile genes were enriched in 424 GO terms (out of 7,277 GO terms tests) and bottom-quartile genes were enriched in 37 GO terms (Supplemental_Dataset_S6). This result raised the concern that the model's strong performance might partly reflect memorization of the characteristics of well-annotated gene families. Ninety-six percent of top-quartile genes and forty-seven percent of bottom-quartile genes were assigned at least one GO term. However, 157 genes in the top decile of predicted probabilities of being phenotype-associated encode proteins with domains of unknown function. Manual curation of the top 15 uncharacterized gene models showed that four genes had only domains of unknown function. These results indicate that the model captured biologically important patterns to distinguish gene models highly likely to be associated with phenotype, and that genes predicted to be associated with phenotypes extend beyond those of well-characterized gene families.

DISCUSSION

A substantial proportion of gene models remain uncharacterized (Pena-Castillo and Hughes, 2007; Schnable and Freeling, 2011; Stoecker et al., 2018). Importantly, this lack of functional characterization is not limited to individual gene models, as some entire gene families also remain poorly understood. The time and resources required to functionally characterize an individual gene remain costly. As a result, the social structure of research and funding allocation decisions incentivize focusing on well-studied genes. Therefore, more than three-fourths of gene models across humans and plants are understudied and have not been functionally characterized (Schnable, 2020; Stoecker et al., 2018). Although the initial impulse may be to assume that all currently uncharacterized gene models are equally good candidates for characterization, here, we demonstrated that it is possible to train models on relatively simple and generalizable sets of features that increase the odds of successful reverse genetics characterization. We curated an

1 extended list of phenotype-associated genes identified in maize from the 1900s to November 27,
2 2021, and identified a complementary set of premature stop enriched-common gene models that
3 can serve as a negative set for model training. All of the genes in this set had evidence of
4 transcription (e.g. expression in RNA-seq studies), many had evidence of translation from
5 proteomic investigations, and more than one-third were conserved at syntenic locations in the
6 sorghum genome, indicating these loss-of-function tolerant genes did not simply represent
7 artifacts produced by *ab initio* gene annotation software. Consistent with this, even after
8 restricting the analysis to gene models classified as true genes by reelGeneScore (threshold > 0.5 ,
9 removing more than one-quarter of maize gene models; Schulz et al., 2025), the evolutionarily
10 informed model still distinguished PAGs with an ROC-AUC of 0.78. This is a decline relative to
11 the ROC-AUC of 0.87 achieved when all annotated genes were included, but still much higher
12 than the 0.5 expected of an uninformative model. Together, these findings support that the
13 predictions generated by the model represent biologically meaningful differences rather than
14 solely the identification of annotation artifacts. The difference between PAGs and PSE-Cs could
15 not be explained by difference in gene family size to which these gene models belong to (p-
16 value=0.13, two-tailed Mann-Whitney *U* test). This was previously found to be an important
17 difference between genes associated with lethal phenotypes and genes associated with non-lethal
18 phenotypes, where presence of even a single paralog reduces the likelihood of lethality under
19 disruption of gene (Lloyd et al., 2015). A simple classifier model provided with evolutionary
20 information and genomic sequence-derived features was able to classify gene models in maize
21 into the correct classes with an ROC-AUC of 0.87.

22 The top quartile of genes predicted to be linked to phenotypes included more than 20% of
23 all maize gene models encoding proteins containing domains of unknown function, more than

300 genes in the top quartile lacked GO annotations, and more than 200 genes lacked any functional description at all, including automated domain annotations. This suggests that our model is not simply learning the properties of well-characterized gene families as uncharacterized genes and members of uncharacterized gene families were not excluded from those predicted to be associated with phenotypes. Syntenic conservation in a closely related species (specifically sorghum for the maize model) consistently ranked as the most important across all trained models. However, when nonsyntenic genes were excluded from both the training and testing set (removing the information content of synteny), a retrained top-100-feature model was able to classify maize gene models with an ROC-AUC of 0.85 compared with 0.87 for the original model, indicating that although synteny is an important predictor, its inclusion is not essential to achieve substantial prediction performance.

Phenotype-associated genes were defined based on reports in the literature describing visible phenotypes in relatively easy-to-characterize tissues. This definition necessarily biased our training dataset against genes whose phenotypes are most apparent in tissues or organs that are more difficult to observe, such as roots. Our initial maize PAG set was also enriched for genes associated with easily characterized kernel phenotypes. Of the 27 newly validated gene models, nearly half were kernel mutant phenotypes. However, kernel-related and non-kernel-related genes showed no significant difference in mean predicted probability of being phenotype-associated (p -value=0.39, two-tailed Mann-Whitney U test), and in fact the newly characterized genes with non-kernel related phenotypes were assigned higher mean probabilities by the model (mean=0.63 vs mean = 0.72) suggesting the high representation of these genes did not bias the model towards genes with phenotypes observed specifically in kernels. Our criteria for the training dataset also specifically excluded genes linked to phenotypes only under specific

1 environmental conditions (e.g. stress specific phenotypes). However, *Arabidopsis* genes
2 associated with conditional phenotypes were not assigned significantly lower predicted
3 probabilities of being phenotype-associated than *Arabidopsis* genes associated with constitutive
4 phenotypes (Supplemental Figure S3).

5 We were able to assess overfitting or species-specific prediction by testing models trained
6 in maize on known phenotype-associated genes from other species. Models trained in maize
7 generalized well across species, assigning substantially higher predicted phenotype probabilities
8 to genes with known mutant phenotypes in both rice and *Arabidopsis*, two distantly related
9 species with large differences in genome size and gene content. Using a top-decile threshold
10 yields very low false positive rates but misses many true phenotype-associated genes, indicating
11 a substantial false negative rate. In the absence of confirmed true negative gene set, predicted
12 probabilities should be interpreted as relative prioritization scores rather than absolute
13 probabilities of phenotype association. It should also be noted that the probability threshold
14 corresponding to any given percentile of scores varied across the species we evaluated.
15 Application of the model may be more comparable when applying similar percentile based
16 threshold than comparable absolute scores. Selection of specific threshold would necessarily be
17 dependent on a user's relative tolerance for false positives and false negatives. As additional
18 genes continue to be cloned and functionally characterized, it will be informative to reassess
19 model performance and determine how well the current patterns hold in larger and more diverse
20 validation sets.

21 One potential application of our classifier is predicting the likelihood that disrupting a
22 given gene will be associated with phenotypic changes. Association studies with genetic markers
23 narrow down the genomic region associated with a given trait of interest. There are cases where

the causal gene is the nearest one to the associated genetic marker (Sahay et al., 2023), especially when there is a high-density genetic marker across the genome and strong linkage disequilibrium decay (Grzybowski et al., 2023). However, for many plants, markers may be at low density, or there may be strong linkage disequilibrium, or in some cases, a combination of both. Having a complementary tool that provides a quantitative score for each gene model indicating the probability that a gene model is linked to a phenotype could guide researchers in prioritizing candidate genes in such association studies. Such a framework could help reduce uncertainty when large genomic intervals contain dozens or even hundreds of annotated gene models. In addition, this tool could provide complementary evidence to traditional QTL mapping or GWAS approaches, ultimately accelerating functional validation. The vast majority of gene models, even in the most studied plants such as maize, rice, and *Arabidopsis*, are yet to be characterized, let alone many orphan crops, which are equally, if not more important (Shrestha et al., 2023). Editing, validating, and characterizing gene models takes a long time and is costly (Stoecker et al., 2018). Quantitative scores for gene models that could be used as confidence scores for prioritizing targets would therefore be invaluable. Such scores could help researchers focus their limited resources on the most promising candidates, thereby accelerating basic discovery and crop improvement. Identifying true genes based on high quantitative scores could mark the starting point for optimizing their functional roles across diverse crop species.

METHODS

Definitions of Phenotype-Associated and Premature Stop Enriched gene sets

The initial set of 272 maize genes whose loss of functional alleles have been experimentally shown to produce a detectable change as an easily observable phenotype, along with the organ affected by each gene, was defined through our manual curation of a larger set of

603 maize locus records annotated as phenotype-associated and retrieved from MaizeGDB (Sen et al., 2023) on November 27, 2021. This was done to exclude loci linked to purely molecular phenotypes (e.g., variation in gene expression) or when the evidence for a genotype-phenotype link came from studies using techniques with higher false-positive rates (e.g., association studies). From the set of approximately 4,500 cloned rice genes in the funRiceGenes database, we manually curated a non-exhaustive set of 100 which met the same phenotype criteria applied to maize (Huang et al., 2022). In four cases, a gene identified in the analysis did not correspond to a gene in the sequence-derived feature set calculated from the MSU v7 rice genome annotations, leaving 96 phenotype-associated rice genes. The set of phenotype-associated genes in *Arabidopsis* was retrieved from previously published list (Lloyd and Meinke, 2012). Excluding gene models in the Lloyd and Meinke list with phenotypes that were classified as ‘NC; non-confirmed’ resulted in a set of 2,018 phenotype-associated genes for *Arabidopsis*. Out-of-sample maize phenotype-associated genes were identified using a set of 2,582 citations (through April 14, 2025) describing maize loci that had not been associated with a phenotype in our original November 27, 2021, dataset. These citations were derived from 682 unique papers, and screening these publications identified an additional 27 loci that met the criteria for publishing a new genotype-to-phenotype association supported by loss-of-function data.

The maize premature stop enriched gene set was defined using a published set of genotype calls derived from whole-genome resequencing of 1,515 maize varieties (Grzybowski et al., 2023, Supplemental_Dataset_S8). The functional consequence of each of the approximately 46 million biallelic genetic variants (SNPs and InDels) identified within this study were annotated using SNPeff v5.1d (De Baets et al., 2012) and the published gene model annotations of the B73_RefGen_V5 reference genome (Hufford et al., 2021). The primary

transcript of each annotated gene model containing a genetic variant that produces a premature stop codon and is present at an alternate allele frequency greater than 0.25 across the population was considered premature stop enriched-common, resulting in a final dataset of 162 premature stop enriched-common maize genes.

Definition of the feature dataset describing maize gene models

A dataset of 14,418 different features describing 39,756 annotated gene models in the B73_RefGen_V5 reference genome was downloaded from the Maize Feature Store in MaizeGDB (Sen et al., 2023, <https://mfs.maizegdb.org/featurebase>). Gene models that were not assigned to any of the ten maize chromosomes (i.e., on unanchored scaffolds) were dropped from the dataset, reducing the total number of gene models to 39,034. Redundant features introduced by the presence of both raw and normalized values in the dataset, features with missing values for more than 1 gene model, annotation of gene model in sorghum version 3.1 annotation and in Tzi8, chromosome assignment, and four subcellular localization features while retaining only one subcellular localization feature, Subcell1 were removed from the dataset. In addition, features describing tri-peptide amino acid composition, which would otherwise have constituted approximately 55% of the total feature set, were also dropped from the dataset. Gene expression and protein abundance values of "none" were converted to "0" prior to analysis. After excluding these features, 2,639 remained, including features describing gene structure, codon usage, nucleotide and protein composition, and non-genomic sequence derived data regarding gene expression atlas, protein expression, chromatin-related features, expression localization, and transcription factor binding sites associated with each gene model in maize. A single non-MaizeGDB-sourced feature, syntenic conservation in sorghum, scored using the dataset published in (Dias et al., 2026), was added, resulting in a final dataset of 2,640 unique features.

Synteny based on sorghum gene models included a filtering step that removed members of tandem duplicate arrays that did not possess reciprocal best BLAST hit relationships between maize and sorghum. The feature generation process and definition of features have been described in detail by Sen et al. (2023). The final feature set of 2,640 features included 31 binary features, 5 categorical features, and 2,604 continuous and discrete features. A total of 350 features were non genomic sequence derived, while the remaining 2,290 features could be directly computed from the maize genome assembly and annotation files.

For use in machine learning applications, categorical features were one-hot encoded, increasing the effective number of variables to 2,675 (Supplemental_Dataset_S9), and continuous and discrete features were min-max normalized across all gene models using the MinMaxScaler function in Scikit-learn v1.3.0 implemented in Python v3.11.4 (Pedregosa et al., 2011).

Random forest classifier model for prediction

Prior to developing or conducting any analysis, approximately 10% of the gene models were set aside as a separate testing dataset. The data were split into training and testing datasets using a previously proposed family-guided strategy (Washburn et al., 2019) and gene family assignments retrieved from PLAZA (Van Bel et al., 2022) to minimize the risk of data leakage. The testing data set included 29 phenotype-associated genes and 19 premature stop enriched-common genes. For each model trained in this study, the remaining 386 labeled gene models were used to optimize the hyperparameters and train the final models. The number of trees in a model, number of features for splitting each tree, maximum depth of the tree, minimum number of samples required to split a node, minimum number of samples required to be at a leaf node, and whether bootstrap samples were used when building trees were optimized via three-fold

cross-validation of the training dataset using the RandomizedSearchCV function in Scikit-learn v1.3.0 implemented in Python (Pedregosa et al., 2011). The optimal set of hyperparameters identified through this process for each model was then used to train the respective final model using the whole training data. ROC curves were plotted based on the model's performance on the hold-out test dataset, and the AUC was used to evaluate the model's performance.

Feature importance scores for both models using 2,675 and 2,290 features were calculated using the built-in feature_importance_ function of a trained random forest model from the Scikit-learn library implemented in Python. This method estimates the importance of each feature based on the mean decrease in impurity across all trees in the ensemble. We were only able to generate values for 99 of the 100 maize features of greatest estimated importance in rice and *Arabidopsis*. One feature, 'gene age' was not included in cross-species prediction. A new model trained on maize data using only those 99 genomic features and the same procedure described above, was employed for cross-species predictions.

Validation of model predictions

Variant density was calculated by dividing the number of variants with minor allele frequency greater than 0.05 from the same resequencing-based dataset used to identify premature stop enriched genes occurring in the CDS of the primary transcript by the CDS length of the primary transcript of that same gene. Additional density values were generated by classifying genetic variants present within the primary transcript of a given gene into low effect, moderate effect, modifier effect, and high effect categories via SnpEff (De Baets et al., 2012). CHH methylation per gene model, rates of synonymous (ks) and non-synonymous substitutions (ka) based on alignment to orthologous genes in sorghum and K_a/K_s data was downloaded from MaizeGDB webpage (Woodhouse et al. 2025, <https://mfs.maizegdb.org/Other>,

<https://mfs.maizegdb.org/DNAmethylation>). The functional impact of genetic markers present between the annotated start and stop sites of each gene of interest was estimated using zero-shot scores provided by PlantCaduceus (Zhai et al., 2025). The absolute value of the log-likelihood difference between reference and alternate alleles was used when the reference or the alternate allele was assigned low relative probabilities, both of which indicate the presence of functional differences segregating in the population. ReelGene scores was downloaded for version 4 maize gene models (Schulz et al., 2025) which were then mapped to version 5 maize gene models. Gene models from version 5 genome annotation that mapped to multiple version 4 gene IDs were excluded resulting in removal of approximately 25% of gene models in version 5 which were likely due to gene annotation error in version 4 but were resolved in version 5. InterProScan V5.75 was used to identify gene models whose translated CDS coded for one or more domains of unknown function (Jones et al., 2014). Tests for GO enrichment or purification were performed using the Python library Goatools v1.4.12 (Klopfenstein et al., 2018). The Gene Ontology annotation file, "Zm00001eb.1.fulldata.txt.gz", for B73 reference genome version 5 was downloaded from maizeGDB (Hufford et al., 2021). To exclude bias introduced by the presence of greater proportions of some gene sets representing entirely uncharacterized proteins, the background set of all analyses was defined as the set of genes associated with one or more GO terms unless otherwise specified. GO term enrichment analysis was calculated as the fold change between the proportion of gene models annotated with a given GO term in the study set and its proportion in the background population set of the gene models.

Curation of rice and *Arabidopsis* gene model features

The rice genome was retrieved from the Rice Genome Annotation Project (Hamilton et al., 2025; Kawahara et al., 2013), and the *Arabidopsis* genome was retrieved from The

Arabidopsis Information Resource (TAIR) (Lamesch et al., 2012). For both species, only data from primary transcripts were used, and genes annotated as transposable elements (TEs) were excluded. Length features including 3' UTR length, 5' UTR length, total exon length, and average exon length were directly calculated from gene annotations in the GFF3 file. *k*-mer ratios were derived using a custom R script that integrated the FASTA sequence with GFF3 annotations. Codon usage frequency, relative synonymous codon usage (RSCU), and GC content at the third codon position (GC3s) were also computed from representative CDS.

Multiple protein sequence descriptors were extracted from one representative peptide sequence per gene using the *protr* R package (Xiao et al., 2015; R Core Team., 2021). After removing stop codons and filtering out sequences shorter than three amino acids, the following descriptor types were computed: triad composition, Geary, Moran, and Moreau-Broto autocorrelations, dipeptide composition, pseudo-amino acid composition (PAAC), quasi-sequence order (QSO), and composition-transition-distribution (CTD: CTDC, CTDD, CTDT). Custom functions with built-in error handling were applied to extract specific subsets of features.

All computations were implemented in parallel to accelerate processing. Syntenic rice genes were identified based on orthology with sorghum genes and directly obtained from the supplementary materials of a previously published study (Dias et al., 2026). Syntenic *Arabidopsis* genes were inferred via comparison with *Brassica rapa* genes using the latest reference genome and the SynOrths pipeline (Cheng et al., 2012). A single feature 'gene age' was not included, as no reliable public resources were found, and they could not be readily derived from sequence or annotation data using straightforward computational methods.

Code availability

All the code used to train, test and evaluate the model, along with the models trained in this study, is available in supplemental code and also in GitHub repository (https://github.com/NikeeShrestha/GenePrediction_Paper).

Acknowledgements

The authors would like to thank scientists from MaizeGDB, especially Dr. Maggie Woodhouse and Dr. Eddy Cannon, for providing information, access to the data used in this study, and help with curating the new phenotype-associated gene models. The authors would also like to thank Jingjing Zhai for providing access to zero-shot scores for maize variants published in MaizeGDB. This research was supported in part by the U.S. Department of Agriculture, National Institute of Food and Agriculture, under award numbers 2020-67021-31528 and 2020-68013-30934, and the AI Research Institutes program supported by NSF and USDA-NIFA under the AI Institute for Resilient Agriculture, Award No. 2021-67021-35329. Zhongjie Ji was additionally supported by a Heuermann Postdoctoral Fellow award from the University of Nebraska–Lincoln.

Author contributions statement

J.C.S. conceived of the project. N.S. processed and analyzed the data. Z.J. generated data for validation of the model on rice and *Arabidopsis*. J.C.S. provided supervision, guidance, and feedback on the design of analyses and approaches. N.S. drafted the paper with input from J.C.S. All authors contributed to the revision of the final manuscript.

Conflict of Interest:

James C. Schnable has equity interests in Data2Bio, LLC; and Dryland Genetics LLC and has performed paid work for Alphabet. He is a member of the scientific advisory board Aflo Sensors. The authors declare no other competing interests.

References

- 1 Alexandrov, N. N., Brover, V. V., Freidin, S., Troukhan, M. E., Tatarinova, T. V., Zhang, H., ...
2 & Feldmann, K. A. (2009). Insights into corn genes derived from large-scale cDNA
3 sequencing. *Plant molecular biology*, 69(1): 179-194.
- 4 Beder, T., Aromolaran, O., Dönitz, J., Tapanelli, S., Adedeji, E.O., Adebiyi, E., Bucher, G. and
5 Koenig, R. (2021). Identifying essential genes across eukaryotes by machine learning.
6 *NAR genomics and bioinformatics*, 3(4):p.lqab110.
- 7 Cheng, F., Wu, J., Fang, L., and Wang, X. (2012). Syntenic gene analysis between brassica rapa
8 and other brassicaceae species. *Frontiers in plant science*, 3:198.
- 9 De Baets, G., Van Durme, J., Reumers, J., Maurer-Stroh, S., Vanhee, P., Dopazo, J.,
10 Schymkowitz, J., and Rousseau, F. (2012). Snpeff 4.0: on-line prediction of molecular
11 and structural effects of protein-coding variants. *Nucleic acids research*, 40(D1):D935–
12 D939.
- 13 Dias, H. M., Sagawa, G. Y., Torres-Rodriguez, V. J., Mural, R. V., Schnable, J. C., and Van
14 Sluys, M.-A. (2026). A sorghum-anchored pan-grass syntenic gene set in grasses. *NAR*
15 *Genomics and Bioinformatics*, 8(2):Iqag030.
- 16 Doebley, J., Stec, A., Wendel, J., and Edwards, M. (1990). Genetic and morphological analysis
17 of a maize-teosinte f2 population: implications for the origin of maize. *Proceedings of the*
18 *National Academy of Sciences*, 87(24):9888–9892.
- 19 Doebley, J., Stec, A., and Hubbard, L. (1997). The evolution of apical dominance in maize.
20 *Nature*, 386(6624):485–488.
- 21 Grzybowski, M. W., Mural, R. V., Xu, G., Turkus, J., Yang, J., and Schnable, J. C. (2023). A
22 common resequencing-based genetic marker data set for global maize diversity. *The*
23 *Plant Journal*, 113(6):1109–1121.

- 1 Guan, J., Wang, J., Niu, W., Peng, Z., Wang, S., Liu, Z., Zhou, G., and Ren, B. (2025). Towards
2 recognizing food types for unseen subjects. *ACM Transactions on Computing for*
3 *Healthcare*, 6(1):1–21.
- 4 Hamilton, J. P., Li, C., and Buell, C. R. (2025). The rice genome annotation project: an updated
5 database for mining the rice genome. *Nucleic Acids Research*, 53(D1):D1614–D1622.
- 6 Huang, F., Jiang, Y., Chen, T., Li, H., Fu, M., Wang, Y., Xu, Y., Li, Y., Zhou, Z., Jia, L., et al.
7 (2022). New data and new features of the funricegenes (functionally characterized rice
8 genes) database: 2021 update. *Rice*, 15(1):23.
- 9 Hufford, M. B., Seetharam, A. S., Woodhouse, M. R., Chougule, K. M., Ou, S., Liu, J., Ricci, W.
10 A., Guo, T., Olson, A., Qiu, Y., et al. (2021). De novo assembly, annotation, and
11 comparative analysis of 26 diverse maize genomes. *Science*, 373(6555):655–662.
- 12 Igartua, E., Contreras-Moreira, B., and Casas, A. M. (2020). Tb1: from domestication gene to
13 tool for many trades. *Journal of Experimental Botany*, 71(16):4621–4624.
- 14 Jones, P., Binns, D., Chang, H.-Y., Fraser, M., Li, W., McAnulla, C., McWilliam, H., Maslen, J.,
15 Mitchell, A., Nuka, G., et al. (2014). Interproscan 5: genome-scale protein function
16 classification. *Bioinformatics*, 30(9):1236–1240.
- 17 Kawahara, Y., de la Bastide, M., Hamilton, J. P., Kanamori, H., McCombie, W. R., Ouyang, S.,
18 Schwartz, D. C., Tanaka, T., Wu, J., Zhou, S., et al. (2013). Improvement of the oryza
19 sativa nipponbare reference genome using next generation sequence and optical map
20 data. *Rice*, 6(1):4.
- 21 Klopfenstein, D. V., Zhang, L., Pedersen, B. S., Ramírez, F., Warwick Vesztröcy, A., Naldi, A.,
22 Mungall, C. J., Yunes, J. M., Botvinnik, O., Weigel, M., et al. (2018). Goatools: A
23 python library for gene ontology analyses. *Scientific reports*, 8(1):10872.

- 1 Kumar, S., Stecher, G., Suleski, M., and Hedges, S. B. (2017). Timetree: a resource for
2 timelines, time, trees, and divergence times. *Molecular biology and evolution*,
3 34(7):1812–1819.
- 4 Kustatscher, G., Collins, T., Gingras, A.-C., Guo, T., Hermjakob, H., Ideker, T., Lilley, K. S.,
5 Lundberg, E., Marcotte, E. M., Ralser, M., et al. (2022a). An open invitation to the
6 understudied proteins initiative. *Nature Biotechnology*, 40(6):815–817.
- 7 Kustatscher, G., Collins, T., Gingras, A.-C., Guo, T., Hermjakob, H., Ideker, T., Lilley, K. S.,
8 Lundberg, E., Marcotte, E. M., Ralser, M., et al. (2022b). Understudied proteins:
9 opportunities and challenges for functional proteomics. *Nature Methods*, 19(7):774–779.
- 10 Kwon, J. J., Pan, J., Gonzalez, G., Hahn, W. C., and Zitnik, M. (2024). On knowing a gene: A
11 distributional hypothesis of gene function. *Cell Systems*, 15(6):488–496.
- 12 Lamesch, P., Berardini, T. Z., Li, D., Swarbreck, D., Wilks, C., Sasidharan, R., Muller, R.,
13 Dreher, K., Alexander, D. L., Garcia-Hernandez, M., et al. (2012). The arabidopsis
14 information resource (tair): improved gene annotation and new tools. *Nucleic acids*
15 *research*, 40(D1):D1202–D1210.
- 16 Lee, I., Ambaru, B., Thakkar, P., Marcotte, E.M. and Rhee, S.Y. (2010). Rational association of
17 genes with traits using a genome-scale gene network for *Arabidopsis thaliana*. *Nature*
18 *biotechnology*, 28(2):149-156.
- 19 Lei, X., Shao, H., Tang, Z., Xu, S., and Zhong, D. (2025). Cross-domain remaining useful life
20 prediction under unseen condition via mixed data and domain generalization.
21 *Measurement*, 244:116451.

- Lloyd, J. P., Seddon, A. E., Moghe, G. D., Simenc, M. C., and Shiu, S. H. (2015). Characteristics of plant essential genes allow for within-and between-species prediction of lethal mutant phenotypes. *The Plant Cell*, 27(8), 2133-2147.
- Lloyd, J. and Meinke, D. (2012). A comprehensive dataset of genes with a loss-of-function mutant phenotype in arabidopsis. *Plant physiology*, 158(3):1115–1129.
- Moore, B.M., Wang, P., Fan, P., Leong, B., Schenck, C.A., Lloyd, J.P., Lehti-Shiu, M.D., Last, R.L., Pichersky, E. and Shiu, S.H. (2019). Robust predictions of specialized metabolism genes through machine learning. *Proceedings of the National Academy of Sciences*, 116(6):2344-2353.
- Oellrich, A., Walls, R.L., Cannon, E.K., Cannon, S.B., Cooper, L., Gardiner, J., Gkoutos, G.V., Harper, L., He, M., Hoehndorf, R. and Jaiswal, P. (2015). An ontology approach to comparative phenomics in plants. *Plant methods*, 11(1):10.
- Pedregosa, F., Varoquaux, G., Gramfort, A., Michel, V., Thirion, B., Grisel, O., Blondel, M., Prettenhofer, P., Weiss, R., Dubourg, V. and Vanderplas, J. (2011). Scikit-learn: Machine learning in Python. *the Journal of machine Learning research*, 12:2825-2830.
- Pena-Castillo, L. and Hughes, T. R. (2007). Why are there still over 1000 uncharacterized yeast genes? *Genetics*, 176(1):7–14.
- R Core Team, 2021. R: A language and environment for statistical computing. R foundation for statistical computing, Vienna, Austria.
- Sahay, S., Grzybowski, M., Schnable, J. C., and Głowacka, K. (2023). Genetic control of photoprotection and photosystem ii operating efficiency in plants. *New Phytologist*, 239(3):1068–1082.

- 1 Schnable, J. C. (2020). Genes and gene models, an important distinction. *New Phytologist*,
2 228(1):50–55.
- 3 Schnable, J. C. and Freeling, M. (2011). Genes identified by visible mutant phenotypes show
4 increased bias toward one of two subgenomes of maize. *PloS one*, 6(3):e17855.
- 5 Schulz, A.J., Zhai, J., AuBuchon-Elder, T., Andorf, C.M., El-Walid, M.Z., Ferebee, T.H.,
6 Gilmore, E.H., Hufford, M.B., Johnson, L.C., Kellogg, E.A., La, T., Long, E.,
7 McMorrow, S.J., Miller, Z.R., Portwood, J.L., II, Romay, M.C., Seetharam, A.S., Stitzer,
8 M.C., Woodhouse, M.R., Wrightsman, T., Buckler, E.S., Monier, B. and Hsu, S-K.
9 (2025), Fishing for a reelGene: evaluating gene models with evolution and machine
10 learning. *Plant J*, 123:e70483.
- 11 Sen, S., Woodhouse, M. R., Portwood, J. L., and Andorf, C. M. (2023). Maize feature store: A
12 centralized resource to manage and analyze curated maize multi-omics features for
13 machine learning applications. *Database*, 2023:baad078.
- 14 Shrestha, N., Hu, H., Shrestha, K., and Doust, A. N. (2023). Pearl millet response to drought: A
15 review. *Frontiers in Plant Science*, 14:1059574.
- 16 Sinha, S., Eisenhaber, B., Jensen, L.J., Kalbuaji, B. and Eisenhaber, F. (2018). Darkness in the
17 human gene and protein function space: widely modest or absent illumination by the life
18 science literature and the trend for fewer protein function discoveries since 2000.
19 *Proteomics*, 18(21-22):1800093.
- 20 Stoeger, T., Gerlach, M., Morimoto, R. I., and Nunes Amaral, L. A. (2018). Large-scale
21 investigation of the reasons why potentially important genes are ignored. *PLoS biology*,
22 16(9):e2006643.

- 1 Upadhyaya, S.R., Bayer, P.E., Tay Fernandez, C.G., Petereit, J., Batley, J., Bennamoun, M.,
2 Boussaid, F. and Edwards, D. (2022). Evaluating plant gene models using machine
3 learning. *Plants*, 11(12):1619.
- 4 Van Bel, M., Silvestri, F., Weitz, E. M., Kreft, L., Botzki, A., Coppens, F., and Vandepoele, K.
5 (2022). Plaza 5.0: extending the scope and power of comparative and functional
6 genomics in plants. *Nucleic Acids Research*, 50(D1):D1468–D1474.
- 7 Wang, T., Birsoy, K., Hughes, N.W., Krupczak, K.M., Post, Y., Wei, J.J., Lander, E.S. and
8 Sabatini, D.M. (2015). Identification and characterization of essential genes in the human
9 genome. *Science*, 350(6264):1096-1101.
- 10 Washburn, J. D., Mejia-Guerra, M. K., Ramstein, G., Kremling, K. A., Valluru, R., Buckler, E.
11 S., and Wang, H. (2019). Evolutionarily informed deep learning methods for predicting
12 relative transcript abundance from dna sequence. *Proceedings of the National Academy*
13 *of Sciences*, 116(12):5542–5549.
- 14 Woodhouse, M. R., Cannon, E. K., Portwood, J. L., Gardiner, J. M., Hayford, R. K., Haley, O.,
15 & Andorf, C. M. (2025). Tools and resources at the maize genetics and genomics
16 database (MaizeGDB). *Cold Spring Harbor Protocols*, 2025(1), pdb-over108430.
- 17 Xiao, N., Cao, D.-S., Zhu, M.-F., and Xu, Q.-S. (2015). protr/protrweb: R package and web
18 server for generating various numerical representation schemes of protein sequences.
19 *Bioinformatics*, 31(11):1857–1859.
- 20 Zhai, J., Gokaslan, A., Schiff, Y., Berthel, A., Liu, Z.-Y., Lai, W.-Y., Miller, Z. R., Scheben, A.,
21 Stitzer, M. C., Romay, M. C., et al. (2025). Cross-species modeling of plant genomes at
22 single-nucleotide resolution using a pretrained dna language model. *Proceedings of the*
23 *National Academy of Sciences*, 122(24):e2421738122.

Figure 1. Differences between premature stop enriched-common genes (PSE-Cs) and

phenotype-associated genes (PAGs). A). Gene models in the B73 reference genome contain alternate alleles that introduce premature stop codons, totaling 713 models. A threshold of 0.25 alternate allele frequency defined 162 premature stop enriched-common genes (PSE-Cs), which were used as a second comparator category of gene models in this study. B). Principal component analysis (PCA) of PSE-Cs ($n=162$) and PAGs ($n=272$). The first two principal components explain 28.62% of the variance. Data included in the PCA included a set of 2,675 previously published features that were downloaded from MaizeGDB (Sen et al., 2023). C). Difference in the number of exons per gene model between PSE-Cs and PAGs. D). Difference in the length of the three prime untranslated regions (3'UTRs) of gene models between PSE-Cs and PAGs. E). Difference in the number of isoforms per gene model between PSE-Cs and PAGs. F). Difference in the proportion of GC content between PSE-Cs and PAGs. ** p-value < 0.01, *** p-value < 0.001, n.s non-significant (two-tailed Mann-Whitney U test).

Figure 2. Model performance using different subsets of total features. A). ROC curves are generated from three models: one using all 2,675 features (non-genomic, genomic sequence and comparative genomics-derived), one using 2,290 genomic sequence and comparative genomics derived features, and one using the top 100 genomic sequence and comparative genomics derived features. Performance was evaluated on 48 withheld gene models. B). Predicted probability from the three models in Panel A across all maize gene models.

Figure 3. Phenotype-associated genes (PAGs) have higher predicted probabilities than uncategorized gene models in maize, rice, and *Arabidopsis*. A). Distribution of predicted probabilities of PAGs, premature stop enriched-common (PSE-C) gene models and

1 uncharacterized maize gene models. **B**). Distribution of predicted probabilities of a subset of
 2 validated gene models ($n=96$) retrieved from the literature in rice (Huang et al., 2022) using a
 3 model trained on maize data with 99 genomic sequence and comparative genomics derived
 4 features. **C**). Distribution of predicted probabilities of a subset of validated gene models
 5 ($n=2,018$) retrieved from the literature in *Arabidopsis* (Lloyd and Meinke, 2012) using a model
 6 trained on maize data with 99 genomic sequence and comparative genomics derived features. **D**).
 7 Predicted probability distribution of newly validated PAGs ($n = 27$, curated from MaizeGDB
 8 since November 27, 2021) compared to hold-out PAGs, hold-out premature stop enriched-
 9 common genes (PSE-Cs), premature stop enriched-rare genes (PSE-Rs) and uncategorized gene
 10 models.

11 **Figure 4. Genes with higher estimated effects on plant phenotypes are conserved in the**
 12 **maize population.** **A**). Bottom and top-quartile gene models were defined from the predicted
 13 probability distribution derived from the model trained with the top 100 genomic sequence and
 14 comparative genomics derived features. **B**). Bottom-quartile gene models show ~ 4.4 -fold higher
 15 odds of containing premature stop codons compared to top-quartile models ($p\text{-value} < 2.26 \times$
 16 10^{-16} , Fisher's exact test) **C**). Variant frequency per kilobase in coding sequences of the primary
 17 transcript of the gene models, separated by effect type. Bottom-quartile models have higher
 18 frequencies of high effect ($p\text{-value} = 9.03 \times 10^{-33}$) and moderate effect variants ($p\text{-value} = 0$),
 19 whereas top-quartile models have higher frequencies of low effect variants ($p\text{-value} = 0$; two-
 20 tailed Mann-Whitney U test). **D**). Proportional distribution of high, moderate, and low effect
 21 variants between quartiles. **E**). Average absolute zero-shot score for variants estimated from a
 22 pre-trained DNA language model, PlantCaduceus (Zhai et al. 2025), in different parts of the gene

- 1 models between the bottom-quartile gene models and top-quartile gene models from panel A. **
- 2 $0.01 < \text{p-value} < 1 * 10^{-20}$, *** $\text{p-value} < 1 * 10^{-200}$ (two-tailed Mann-Whitney U test).







