

# Unsupervised discovery of functional sequence patterns from protein language model with MotifAE

Received: 26 February 2026

Chao Hou<sup>1</sup>✉, Di Liu<sup>2</sup> & Yufeng Shen<sup>1,2,3</sup>✉

Accepted: 24 August 2026

Published online: 02 September 2026

 Check for updates

Protein language models (pLMs) learn sequence patterns at evolutionary scale, but these patterns remain inaccessible within these “black box” models. To discover them, we developed MotifAE, an unsupervised framework based on the sparse autoencoder (SAE) architecture that projects pLM embeddings into an interpretable, sparse latent space. MotifAE introduces an additional smoothness loss to encourage coherent feature activation, which markedly improves the identification of known functional motifs compared to the standard SAE. The sequence patterns captured by MotifAE exhibit rich diversity, align with known functional motifs, and are reflected in the model’s weight space. Beyond short motifs, MotifAE also captures some structural domains, with latent feature activation scores correlating with residue importance for diverse domain functions. By aligning MotifAE features with experimental data, we further identified features associated with domain folding stability. These features enable the prediction of a stability-specific fitness landscape. Overall, MotifAE provides a general framework for systematic sequence pattern discovery and interpretation, with the potential to advance protein function analysis, mutation effect interpretation, and rational protein engineering.

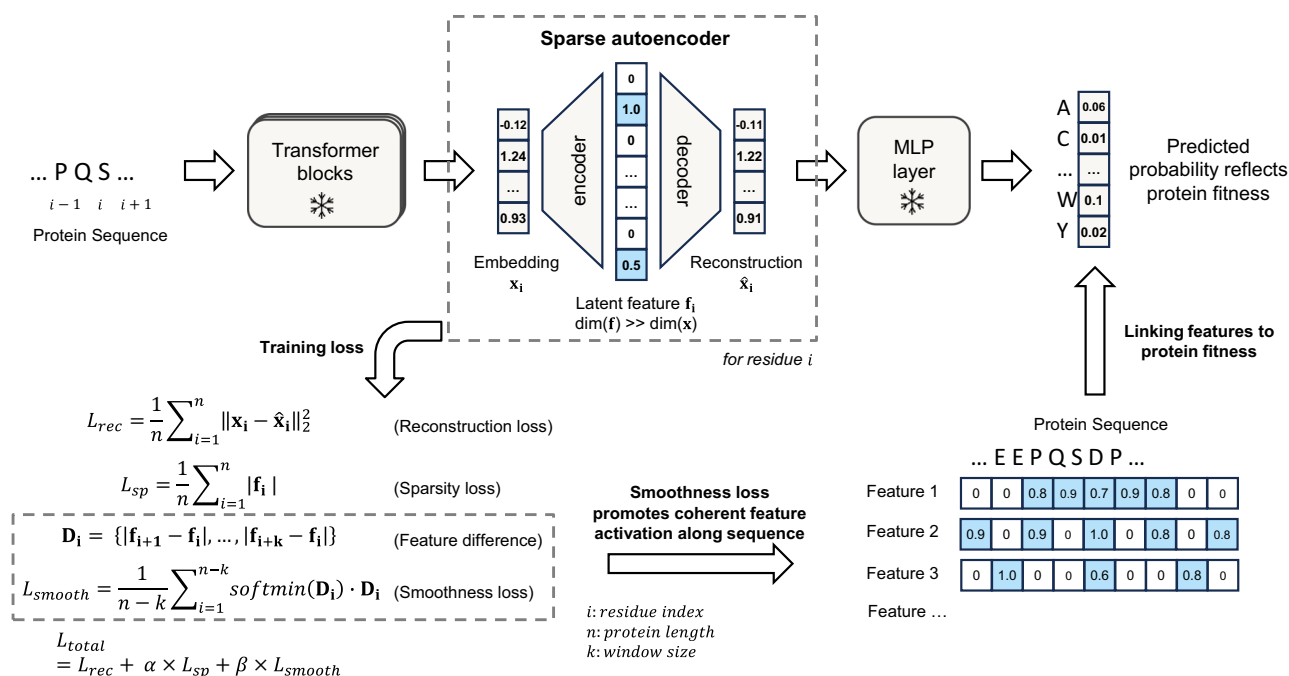
Proteins mediate essential biological functions. Their functional elements, such as catalytic centers, protein/nucleic acid/small molecule binding interfaces, and post-translational modification sites, exhibit conserved sequence patterns commonly referred to as motifs. Experimental identification of a motif requires first locating protein regions that share the same function and then aligning them to get the sequence pattern, which is both labor-intensive and time-consuming. As of 2026, only ~350 motifs have been curated in the Eukaryotic Linear Motif (ELM) database<sup>1</sup>. To complement experimental efforts, machine learning methods have been developed<sup>2–4</sup>, often trained on curated resources like ELM. However, these supervised models are typically tailored to specific functions and inherently biased toward their training datasets, limiting their generalizability. Therefore, systematic

and unbiased discovery of functional motifs is crucial for advancing our understanding of protein functions.

Besides short motifs, structural domains also exhibit conserved sequence patterns. By aligning experimentally solved structures in the Protein Data Bank<sup>5</sup> (PDB) and high-quality predicted structures, domain sequence patterns have been systematically characterized<sup>6,7</sup>. However, the functional organization within domains remains poorly understood. In particular, it is unclear how individual residues contribute to folding stability, catalytic activity, partner binding, or other functional roles. Characterizing residue-level importance across these functions is essential for understanding domain organization and has the potential to enable the rational engineering of domains with desired properties.

<sup>1</sup>Department of Systems Biology, Columbia University Irving Medical Center, New York, NY, USA. <sup>2</sup>Department of Biomedical Informatics, Columbia University Irving Medical Center, New York, NY, USA. <sup>3</sup>JP Sulzberger Columbia Genome Center, Columbia University, New York, NY, USA.

✉ e-mail: [ch3849@cumc.columbia.edu](mailto:ch3849@cumc.columbia.edu); [ys2411@cumc.columbia.edu](mailto:ys2411@cumc.columbia.edu)



**Fig. 1 | The MotifAE framework.** MotifAE uses the sparse autoencoder (SAE) architecture with an additional smoothness loss. MotifAE projects ESM2-650M embeddings (dimension 1280) into a high-dimensional latent space (dimension 40,960) using an encoder, then reconstructs the embeddings through a decoder. The reconstructed embeddings can be projected to amino acid probabilities via the

fixed ESM2 MLP layer, enabling prediction of the protein fitness landscape. MotifAE is trained with three objectives: a reconstruction loss, a sparsity loss, and a smoothness loss, the latter encourages coherent latent feature activations. Three MotifAE models with smoothness loss window sizes of three, four, and five were trained.

In recent years, deep learning has driven remarkable advances in modeling protein sequence<sup>8</sup>, structure<sup>9</sup>, and function<sup>10</sup>. A major development is protein language models (pLMs), such as ESM2<sup>8</sup>, which are self-supervised models trained on large-scale protein databases using masked or next-token prediction. Through this pre-training, pLMs implicitly learn conserved sequence patterns at evolutionary scale. For example, studies have shown that pLM predicts protein structure by capturing pairwise contact motifs<sup>11</sup>, pLM attentions identify structural contacts and target binding sites<sup>12</sup>, and pLM embeddings show similarity for proteins within the same domain family<sup>13</sup>. Moreover, numerous studies have leveraged pLM embeddings to predict functional regions such as protein sorting signals, transmembrane regions and post-translational modification sites<sup>3</sup>. Collectively, these studies demonstrate that pLMs encode rich information about the sequence patterns of functional motifs and structural domains. However, because pLMs operate as “black boxes”, the sequence patterns they encode are not directly accessible.

Various attempts have been made to extract interpretable features from language models. Among them, sparse autoencoders (SAEs) have shown strong potential<sup>14</sup>. The core assumption of SAE is that the model embedding can be expressed as a linear combination of a set of different features<sup>15</sup>, with each feature ideally corresponding to an interpretable underlying factor. To achieve this, SAE uses an encoder to project embeddings into a higher-dimensional, sparse latent space, and uses a decoder to reconstruct the embeddings from this space. During SAE training, a reconstruction loss is used to ensure that the embeddings are accurately reconstructed, and a sparsity constraint is used to encourage sparse feature activations. Hereafter, we refer to each neuron in the SAE latent space as a feature, and neurons with positive values are activated. Typically, only tens of features are active per token, far fewer than the original embedding dimension, which aids interpretability. SAEs have been successfully applied in large language models (LLMs) to find interpretable semantic

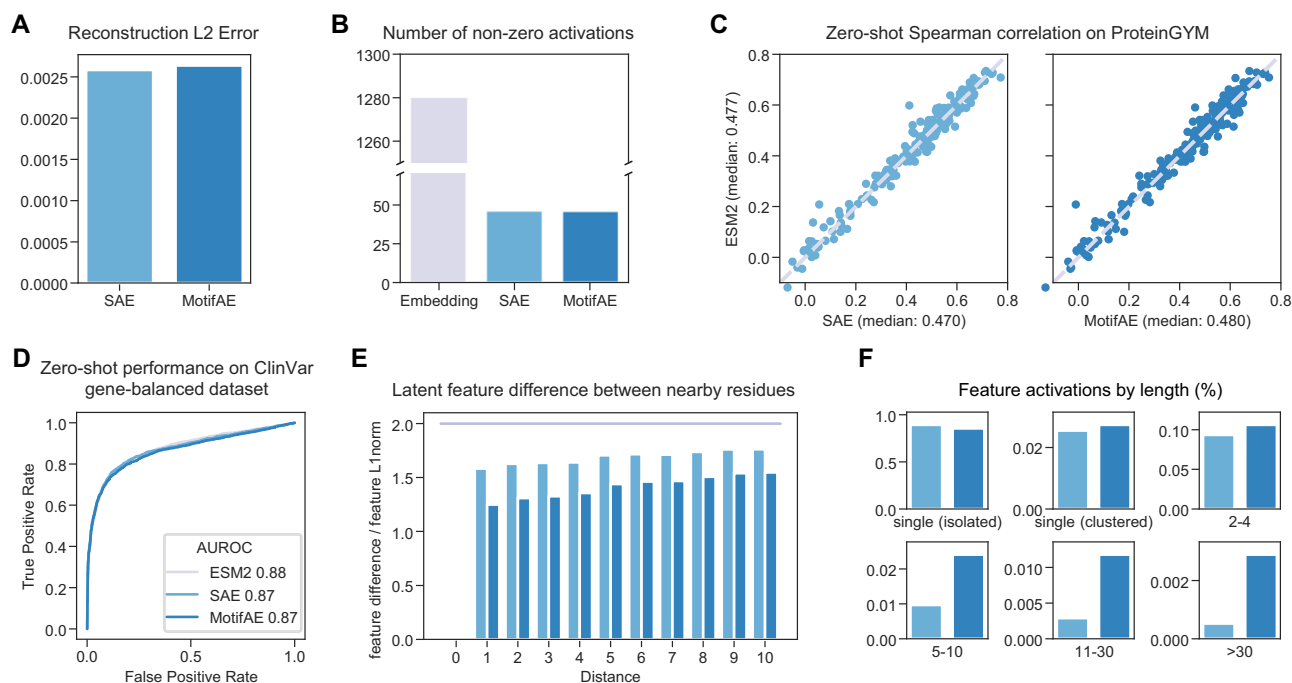
features<sup>14,15</sup>. Adapting the training and interpretation strategies developed for LLM SAE, recent works found that applying SAE to pLMs has the potential to identify functional regions and structural elements<sup>16–18</sup>. However, it remains unclear whether biologically informed priors can improve pLM SAE interpretability and whether SAEs can systematically reveal motif sequence patterns and domain functional organization.

Here, we develop MotifAE, introducing methodological advances in both SAE training and interpretation to better suit it to biological sequences and to experimental data. We incorporate an additional smoothness loss to encourage MotifAE’s latent features to activate coherently, which markedly improves the identification of known functional motifs compared to the standard SAE. To analyze the motif sequence patterns captured, we compute position-specific scoring matrices (PSSMs) for MotifAE features, which reveals rich sequence diversity and highlights MotifAE’s potential for systematic functional motif discovery. For longer sequence patterns, MotifAE significantly outperforms the standard SAE in domain identification, with latent feature activation scores correlating with residue importance for different domain functions. Furthermore, by supervised alignment with experimental data, we identify MotifAE features associated with domain folding stability; these features subsequently improve performance on domain stability prediction and support stability-guided sequence redesign, as evaluated *in silico*. Overall, MotifAE provides a framework for discovering functional sequence patterns from pLMs and annotating them using curated databases and experimental data.

## Results

### MotifAE is a sparse autoencoder with coherent latent feature activation

We developed MotifAE, an adaptation of the sparse autoencoder (SAE) for discovering functional sequence patterns from pLMs. A standard SAE is trained with a reconstruction loss and a sparsity constraint (Fig. 1), ensuring that the original embeddings are accurately



**Fig. 2 | Comparison of MotifAE and SAE in terms of reconstruction error, latent sparsity, fitness prediction performance, and activation length distribution.**

**A** ESM2 embedding reconstruction error on the evaluation set (see Methods), averaged across embedding dimensions. **B** Mean number of non-zero neurons per residue. The ESM2 embedding has 1280 non-zero neurons, while both SAE and MotifAE have ~46 non-zero (positive) latent features per residue. **C** Zero-shot performance (Spearman correlation) on ProteinGYM DMS datasets, only single substitution mutations were evaluated. Each dot represents a DMS dataset. Wild-type marginal log-likelihood ratio was used to estimate mutation effects. **D** Zero-shot

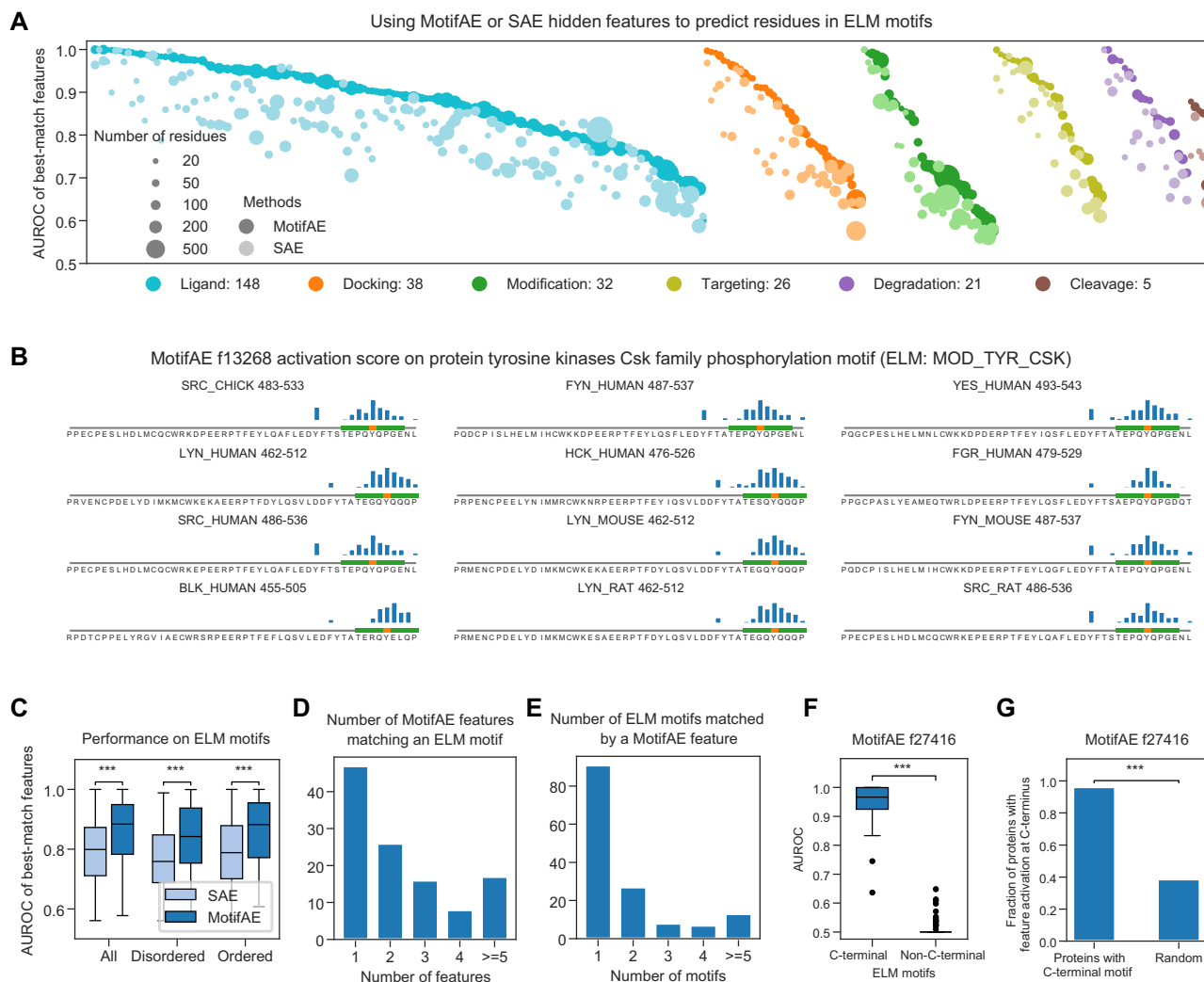
performance on classifying pathogenic and benign missense mutations in the ClinVar gene-balanced dataset (see Methods). **E** Latent feature difference between nearby residues, measured using the absolute activation difference divided by the L1 norm of the activations. The value on y-axis is 0 for identical activations and -2 for random sparse activations. **F** Length distribution of activations from all features. For each feature, a single (clustered) activation is defined as an activated single residue with other activated residues within two amino acids along the sequence. In all plots, light blue denotes SAE, and dark blue denotes MotifAE. Source data are provided as a Source Data file.

reconstructed while enforcing sparse activation of latent features. To enhance biological relevance, we introduced a smoothness loss (see Methods), which requires that at least one neighboring residue within a defined window exhibits latent feature activations similar to the center residue (Fig. 1), thereby promoting coherent feature activation along the sequence. The neighboring residues with similar activations are not necessarily directly adjacent. This design is motivated by the observations that protein motifs often involve contiguous residues, that basic structural elements exhibit characteristic local patterns, such as alternating interactions in  $\beta$ -strands and periodic interactions in  $\alpha$ -helices, and that structural domains are assembled from functional motifs and basic structural elements. We trained MotifAE models with smoothness loss window sizes of three, four, and five residues. We did not explore larger window sizes as they would weaken the smoothness constraint. Because the smoothness loss is computed over the entire sequence, it can propagate along the sequence to capture some longer-range relationships. MotifAE models with different window sizes show similar performances across diverse evaluations. We focus on the model with window size of three in main figures and summarize results for all models in Table S1.

We trained MotifAE on the final-layer embeddings from ESM2-650M, which has shown strong performances across diverse downstream tasks. ESM2-650M captures the protein fitness landscape and is widely used to predict mutation effects<sup>19–21</sup>. We used the last-layer embeddings because they are directly used to predict amino acid probability, allowing us to investigate the relationship between latent features and fitness<sup>19,20</sup> (Fig. 1). To ensure that MotifAE captures representative protein features, training was performed on 2.3 million representative proteins obtained from structure-based clustering of the AlphaFold2 structure database<sup>22</sup>. The dimension of the latent space

and weights for different losses were determined by jointly considering reconstruction error, latent sparsity, and the protein fitness prediction performance (Fig. S1; see Methods). Based on these criteria, we selected a hidden dimension of 40,960, corresponding to a 32-fold expansion over the ESM2 embedding dimension. Notably, we found that the reconstruction of ESM2 embeddings is highly sensitive: although the reconstruction errors show little variation across models with different latent space dimensions (Fig. S1B), the quality of the reconstructed embeddings varies substantially, as reflected in their performance on protein fitness prediction (Fig. S1D, E).

We trained a standard SAE without the smoothness loss for comparison (see Methods). Both MotifAE and SAE achieve low reconstruction errors (Fig. 2A), with approximately 46 latent features activated per residue on average (Fig. 2B). To assess the quality of the reconstructed embeddings, we evaluated their performance on protein fitness prediction. We projected the reconstructed embeddings through the fixed ESM2 MLP layer to predict amino acid probability (Fig. 1), calculated the wild-type marginal log-likelihood ratio (LLR; see Methods)<sup>19,21</sup>, and used the LLR to estimate mutation effects. Across deep mutational scanning (DMS) experiments in ProteinGYM<sup>20</sup> and pathogenic/benign mutations in ClinVar<sup>23</sup> (see Methods), reconstructions from both MotifAE and SAE perform comparably to the original ESM2 embeddings (Fig. 2C, D). The smoothness loss of MotifAE was introduced to encourage coherent latent feature activation. To evaluate this, we first analyzed feature activation differences between neighboring residues and found that MotifAE features are more similar among nearby residues than those from the standard SAE (Fig. 2E). We further analyzed the length distribution of activations and found that MotifAE features more frequently form clustered and contiguous activations than those of SAE (Fig. 2F). Activated regions with a length



**Fig. 3 | MotifAE captures known functional motifs.** **A** AUROC of the best-match feature for each ELM motif. Each column shows the performance of MotifAE (dark dots) and SAE (light dots) for an ELM motif. Dot size indicates the number of residues in verified motif regions, and colors represent different ELM categories. The six categories and the number of motifs in each are listed below the plot. **B** Feature activation along protein sequences. The height of the blue bars indicates the normalized feature activation scores, the green region denotes the motif, and the orange region marks the phosphorylated tyrosine residue. Only the C-terminal 50 residues of these proteins are shown. **C** Distribution of AUROCs of best-match features for ELM motifs. Disordered and ordered regions were classified using IUPred3<sup>36</sup>. Only motifs with at least 20 residues located in disordered or ordered verified regions were included. The box extends from the first quartile to the third quartile, with a line marking the median, whiskers span to the most extreme points

within  $1.5\times$  the interquartile range. Paired two-sided t-tests: All,  $p$ -value =  $1.99e-60$  ( $n = 270$ ); disordered regions,  $p$ -value =  $1.77e-39$  ( $n = 179$ ); and ordered regions,  $p$ -value =  $2.63e-37$  ( $n = 188$ ). \*\*\*,  $p$ -value < 0.001. **D** Number of features matching an ELM motif, and **E** Number of motifs matched by a MotifAE feature; a motif–feature match is defined by AUROC > 0.9. **F** AUROC of MotifAE feature f27416 on C-terminal versus non-C-terminal ELM motifs. Non-C-terminal ELM motifs are defined as those with no verified regions within 10 residues of the C-terminus. Mann–Whitney U test  $p$ -value =  $2.80e-35$  ( $n = 25$  for C-terminal motifs and  $n = 212$  for non-C-terminal motifs). **G** Percentage of proteins with activation of f27416 at the C-terminus. Proteins with known C-terminal motifs were obtained from the ELM dataset, random proteins were sampled from the evaluation set. Fisher’s exact test  $p$ -value =  $8.66e-86$  ( $n = 320$  for proteins with C-terminal motifs and  $n = 1000$  for random proteins). Source data are provided as a Source Data file.

of 5–30 residues suggest possible motifs, while those longer than 30 residues suggest possible domains.

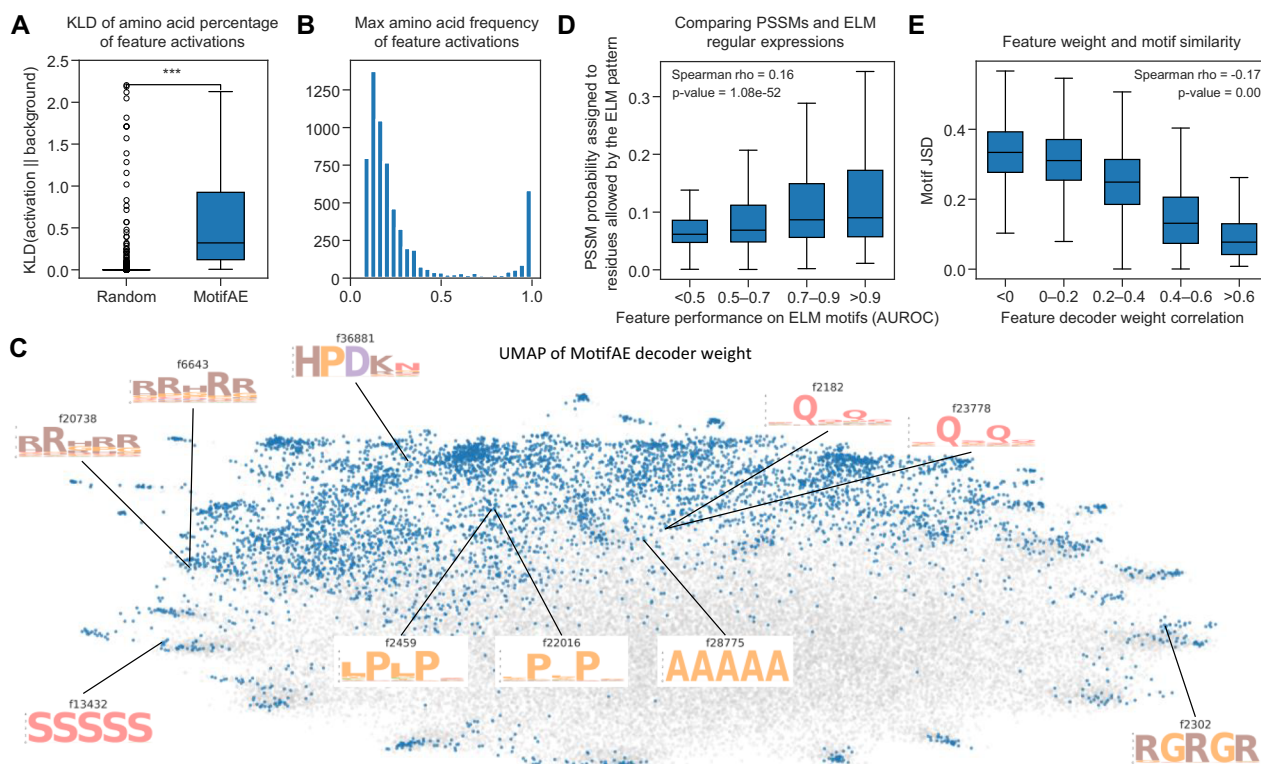
Together, these results demonstrate that MotifAE preserves embedding reconstruction quality and latent space sparsity, while the smoothness loss promotes coherent feature activations, with the potential to enhance the identification of functional motifs and structural domains.

### MotifAE captures known functional motifs

We first investigated whether MotifAE latent features capture functional motifs. Experimentally validated motifs were obtained from the ELM<sup>1</sup> database, which covers six categories: ligand-binding, docking, post-translational modifications, targeting signals, degradation, and proteolytic cleavage sites. Each ELM motif contains multiple

experimentally verified functional regions. We focused on 270 motifs with at least 20 residues in verified regions, as motifs with fewer annotated residues might be matched by chance due to the large number of latent features. We compared all MotifAE features against all ELM motifs, assessing whether the feature activation scores could distinguish residues located inside versus outside the experimentally verified motif regions (see Methods).

Considering the best-match feature for each motif, MotifAE demonstrates strong performance, achieving a median AUROC of 0.88 across the 270 motifs, significantly outperforming the standard SAE (median AUROC 0.80) (Fig. 3A). The best-match MotifAE features achieve AUROCs above 0.8 for 193 motifs and above 0.9 for 114 motifs. To assess statistical significance, we performed permutation tests by shuffling the activation scores of the best-matching features while



**Fig. 4 | The repertoire of motif sequence patterns captured by MotifAE.** **A** KL divergence between the amino acid composition of activated regions of each feature and the background (amino acid composition in the representative protein dataset). Random peptides matched in number and length distribution were used as the control. Mann–Whitney U test  $p$ -value = 0 ( $n = 6421$ ); \*\*\*:  $p$ -value < 0.001. **B** Frequency of the most abundant amino acid in the activated regions of each feature. **C** UMAP visualization of MotifAE decoder weights. Each point represents a feature; gray points correspond to features without coherent activations, while blue points correspond to features with coherent activations. PSSM sequence logos are shown for some features. **D** The x-axis shows the AUROC of using the feature activation score to identify ELM motif regions, and the y-axis shows the PSSM

probability of residues permitted by the corresponding ELM regular expression. A value of 0.05 represents expectation of random matches; higher values indicate greater similarity between the feature PSSM and the ELM regular expression (see Methods). **E** Relationship between feature similarity of decoder weights and similarity of PSSMs. Decoder weight similarity was measured using Pearson correlation, while PSSM similarity was measured using Jensen–Shannon divergence (JSD), where smaller JSD values indicate higher similarity. The box extends from the first quartile to the third quartile, with a line marking the median, whiskers span to the most extreme points within  $1.5\times$  the interquartile range. Source data are provided as a Source Data file.

preserving the activation length distribution (see Methods). This analysis confirms that the observed performance is highly unlikely under the null distribution, and that MotifAE is more significant than SAE (Fig. S2A, B). One representative motif–feature pair is the ELM motif MOD\_TYR\_CSK, the C-terminal phosphorylation motif in Src-family proteins targeted by the non-receptor tyrosine kinase Csk family<sup>24</sup>, with f13268. Across 12 proteins with a known MOD\_TYR\_CSK motif in ELM, f13268 predicts residues in the motif with an AUROC of 0.97. Notably, the highest activation of f13268 occurs at the phosphorylated tyrosine residue (Fig. 3B), indicating that this feature captures the biological characteristics of the motif. Importantly, MotifAE performs consistently across structural contexts, with median AUROCs of 0.88 in ordered regions and 0.84 in disordered regions, compared to 0.79 and 0.76, respectively, for SAE (Fig. 3C; see Methods).

Subsequently, we analyzed the specificity of the relationship between ELM motifs and MotifAE features. We defined a matched motif–feature pair as one with AUROC > 0.9, resulting in 322 matched pairs involving 114 ELM motifs and 146 MotifAE features (Fig. S3). Among these, 47 ELM motifs are matched by a single MotifAE feature, whereas 67 motifs are matched by multiple features (Fig. 3D). On the feature side, 91 MotifAE features match only one motif (Fig. 3E), whereas 55 features match multiple motifs. Notably, among the features that match multiple motifs, f27416 is one of the most prominent: it matches 17 motifs, all located at the C-terminus of proteins. Across all

ELM motifs, f27416 achieves a median AUROC of 0.97 for 25 motifs exclusively located at the C-terminus but shows no predictive power for other motifs (Fig. 3F). In addition, f27416 is not universally activated at the C-terminus of all proteins: among 320 proteins with known C-terminal motifs in ELM, 96% exhibit f27416 activation, compared with only 39% of randomly sampled proteins (Fig. 3G).

Overall, these results demonstrate that the smoothness loss markedly improves MotifAE’s ability to capture functional motifs. These results also show that MotifAE features exhibit varying levels of granularity: some correspond to specific functions, while others capture more general and shared functions.

### The repertoire of motif sequence patterns captured by MotifAE

To characterize the sequence patterns of short motifs captured by MotifAE, we recorded activated regions in the representative protein dataset<sup>22</sup> for each feature, considering only regions with a length of 5–30 residues. Among the 40,960 features, 6,421 display coherent activations. Since functional motifs often exhibit distinct sequence preferences, we analyzed these activated regions and found that they exhibit significantly biased amino acid composition compared to random peptides (Fig. 4A). Notably, approximately 400 features are almost exclusively activated by a single amino acid (Fig. 4B; with one amino acid type accounting for >99% of all activations). To describe the sequence patterns, we constructed position-specific scoring

matrices (PSSMs) for each MotifAE feature: activated regions (5–30 residues) were aligned using GibbsCluster<sup>25</sup>; the alignment cores, set to five residues for all features, were used to calculate the PSSMs (see Methods). The resulting PSSMs exhibit median KL divergence values (averaged across the five residues) relative to the background of 2.1 (Fig. S4A). Among the PSSMs with well-defined sequence patterns, as indicated by higher KL divergence, some correspond to homopolymeric repeats (e.g., f13432 and f28775), while some exhibit multiple amino acid types (e.g., f2302 and f36881) (Fig. 4C).

To evaluate the validity of these PSSMs, we compared them against ELM motifs. Specifically, we aligned each feature PSSM to each ELM regular expression and analyzed the PSSM probabilities of amino acids permitted by the ELM expression (see Methods). We observed that MotifAE features that better identify motif regions also better match the corresponding ELM regular expressions (Fig. 4D). Next, we analyzed the latent representation of each feature. In the MotifAE decoder, each feature has a weight vector that is multiplied by its activation score to reconstruct the embedding. The decoder weights encode the properties of features in the embedding space. We found that features with similar decoder weight vectors tend to produce similar PSSMs (Fig. 4E). For example, features f6643 and f20738 both capture an “RRHRR” pattern; features f2459 and f22016 both capture an “LPLP” pattern; and features f2182 and f23778 both exhibit paired “Q” residues separated by a gap (Fig. 4C). We also observed that features with coherent activations tend to occupy denser neighborhoods in decoder weight space than others (Figure S4B). While we used an alignment core length of five residues when calculating the PSSMs, we tested core lengths of four and six residues and found that results are consistent (Fig. S4A, C, D). Together, by systematically characterizing the motif sequence patterns captured by MotifAE, we found that they exhibit specific sequence preferences, these sequence patterns align with known functional motifs and are reflected in the MotifAE weight space.

### MotifAE captures the functional organization of domains

Next, to investigate whether MotifAE captures some longer sequence patterns, we assessed the ability of its feature activation scores to distinguish residues located inside versus outside annotated structural domains. We evaluated 400 CATH<sup>7</sup> domains, including all major CATH classes: mainly  $\alpha$ -helical, mainly  $\beta$ -sheet, mixed  $\alpha/\beta$ , and others (see Methods). Considering the best-matching feature for each domain, MotifAE significantly outperforms the SAE across all CATH classes (Fig. 5A). A permutation test further confirms the significance of this result (Fig. S2C, D and see Methods). Additionally, we analyzed the specificity of the relationship between CATH domains and MotifAE features. We defined a matched domain–feature pair as one with AUROC > 0.8, resulting in 181 matched pairs involving 105 domains and 106 features. Among these, 71 domains are matched by a single MotifAE feature (Fig. 5B). On the feature side, 84 MotifAE features match only one domain (Fig. 5C). However, performance decreases for domains with high contact order (defined as the mean sequence distance between all contacting residue pairs within a domain) (Fig. S5A, B). This reflects the limitation of MotifAE in capturing complex long-range sequence patterns, which is expected given that the smoothness loss operates over a short sequence window. We note that this analysis is not intended to position MotifAE as a domain identification method, as structure alignment approaches already address that task effectively.

Furthermore, we investigated whether MotifAE captures the functional organization of domains: the residue importance for different domain functions. The mean mutation effect per residue measured by deep mutational scanning (DMS) experiments was used as a proxy of residue importance for functions. We analyzed 177 DMS datasets from ProteinGYM (see Methods), covering measurements of stability, organismal fitness (growth), activity, expression, and binding.

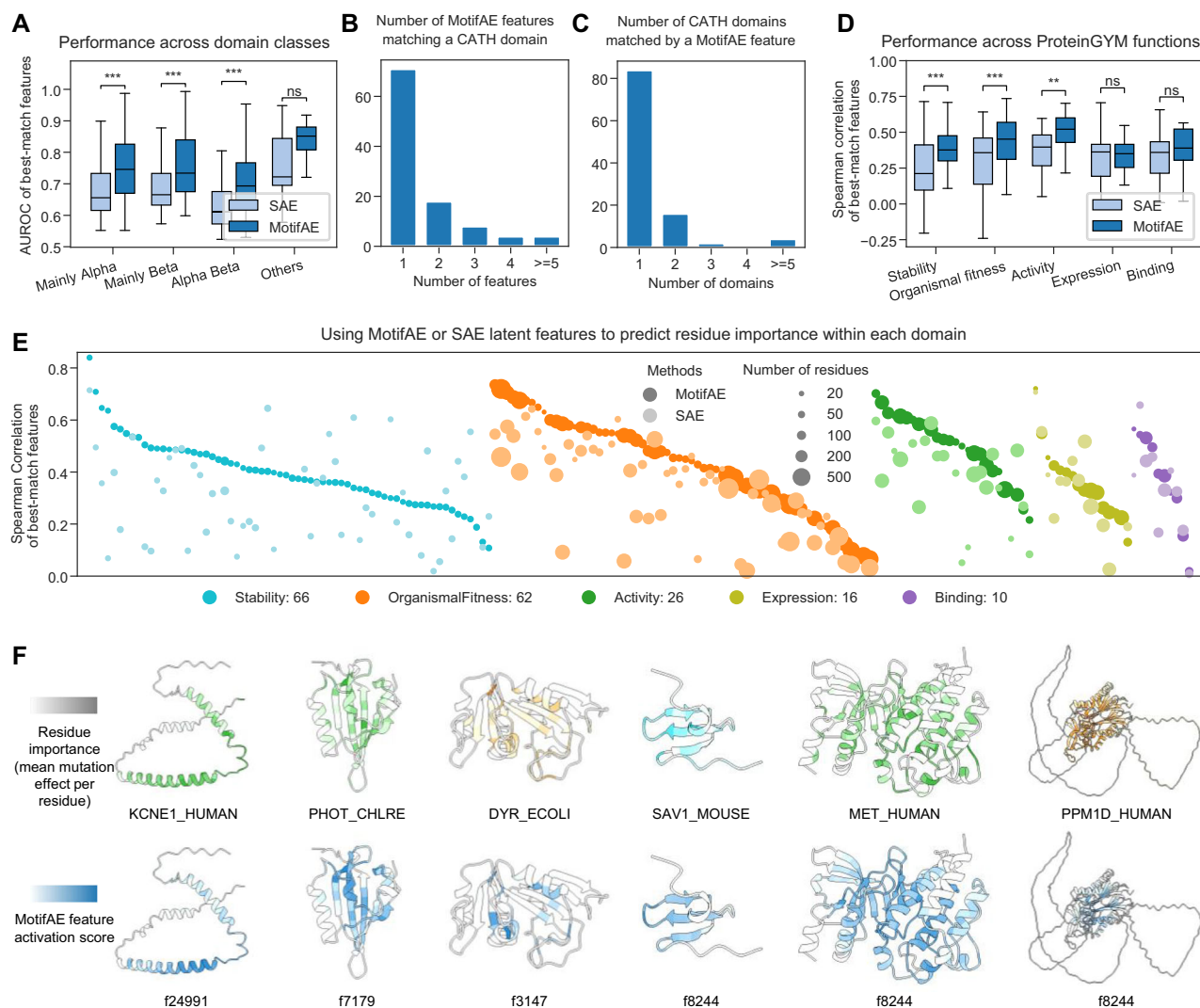
Spearman correlation was used to quantify whether latent feature activation scores capture residue importance. Considering the best-match feature for each experiment, MotifAE achieves a median Spearman correlation of 0.41, significantly outperforming the standard SAE (0.33), with higher correlations observed for 80% of experiments (Fig. 5D, E). The improvement is particularly pronounced for residue importance for stability (median Spearman correlation 0.38 vs. 0.21), activity (0.52 vs. 0.40), and organismal fitness (0.45 vs. 0.36) (Fig. 5D). Among individual features, f8824 is the most prominent, achieving a median Spearman correlation of 0.36 (Fig. S5C), slightly exceeding that of the best-match SAE features for each experiment. Examples of MotifAE features capturing the functional organization of domains are visualized on protein structures (Fig. 5F), showing strong alignment between MotifAE latent feature activation scores and residue importance for different functions.

### Predicting a stability-specific fitness landscape via supervised MotifAE feature selection

pLM-predicted amino acid probabilities capture the protein fitness landscape and are widely used for mutation effect prediction and protein engineering. MotifAE retains essential information about the fitness landscape (Fig. 2C), and its amino acid probabilities are computed from reconstructed embeddings that integrate all latent features weighted by their activation scores (Fig. 1), with different features capturing different functions and properties. Both pLM and MotifAE therefore predict a general fitness landscape that reflects combined constraints from multiple functional and biophysical requirements. Here, we investigated whether we could identify a subset of MotifAE features associated with a specific property, whether these features can be used to adjust predicted amino acid probabilities to reflect a property-specific fitness landscape, and whether this approach can improve prediction of that property and enable the rational engineering of domains with desired property levels.

We developed a supervised approach, called MotifAE-G, to identify features associated with a specific property from DMS experiments (Fig. 6A). MotifAE-G introduces a binary gate to select which features contribute to the embedding reconstruction. The reconstructed embedding, based on the selected features, is then passed through the ESM2 MLP layer to predict amino acid probabilities and compute the LLR. A differentiable Spearman correlation between the LLR and the experimental measurement serves as the loss function. Importantly, this loss is used to update only the binary gate, while all parameters of ESM2 and MotifAE models remain fixed (see Methods). Following training, the binary gate identifies MotifAE features associated with the property of interest, and the residues activated by these features likely underlie the corresponding property. We note that this is a lightly supervised method designed to identify latent features associated with a given property, rather than to maximize predictive performance for that property. When prediction is the goal, specialized predictors or more capable models trained for that task can perform better.

We used the mega-scale domain folding stability dataset<sup>26</sup>, focusing on DMS experiments of 412 domains (see Methods). Both ESM2 and MotifAE perform well on this dataset, with median Spearman correlations of 0.43 and 0.44, respectively (Fig. S6A). To identify stability-associated features, we split the 412 domains into training, validation, and test sets, ensuring less than 30% Foldseek<sup>27</sup> structure token identity between them (see Methods). Using the MotifAE-G framework, we selected 1,583 stability-associated features from the training data. With these features, MotifAE-G achieves median Spearman correlations of 0.60, 0.54, and 0.55 for the training, validation, and test sets, respectively (Fig. S6B), which is comparable to that of FoldX<sup>28</sup>, a biophysics-based model. We note that using a continuous gate, or training more of ESM2's parameters, can lead to better performance (Fig. S6B, see Methods). To assess the biophysical relevance of selected features, we first evaluated how they capture residue



**Fig. 5 | MotifAE captures the functional organization of domains.** **A** Distribution of AUROCs for the best-match feature corresponding to each CATH domain, stratified by domain class. The box extends from the first quartile to the third quartile, with a line marking the median, whiskers span to the most extreme points within 1.5× the interquartile range. Paired two-sided t-test was used. Mainly Alpha:  $p$ -value =  $1.21\text{e} - 14$  ( $n = 95$ ); Mainly Beta:  $p$ -value =  $1.55\text{e} - 10$  ( $n = 77$ ); Alpha Beta:  $p$ -value =  $2.60\text{e} - 38$  ( $n = 220$ ); Others:  $p$ -value =  $0.068$  ( $n = 8$ ). \*\*\*,  $p$ -value < 0.001; ns: not significant. **B** Number of MotifAE features matching each CATH domain, and **C** Number of CATH domains matched by each MotifAE feature; a domain–feature match is defined as AUROC > 0.8. **D** Performance of the best-match features for each DMS experiment, stratified by ProteinGYM experimental categories. Performance is quantified using the Spearman correlation between the feature activation score and the mean mutation effect per residue, which serves as a proxy for residue importance across different functions. Paired two-sided t-test was used. Stability:  $p$ -

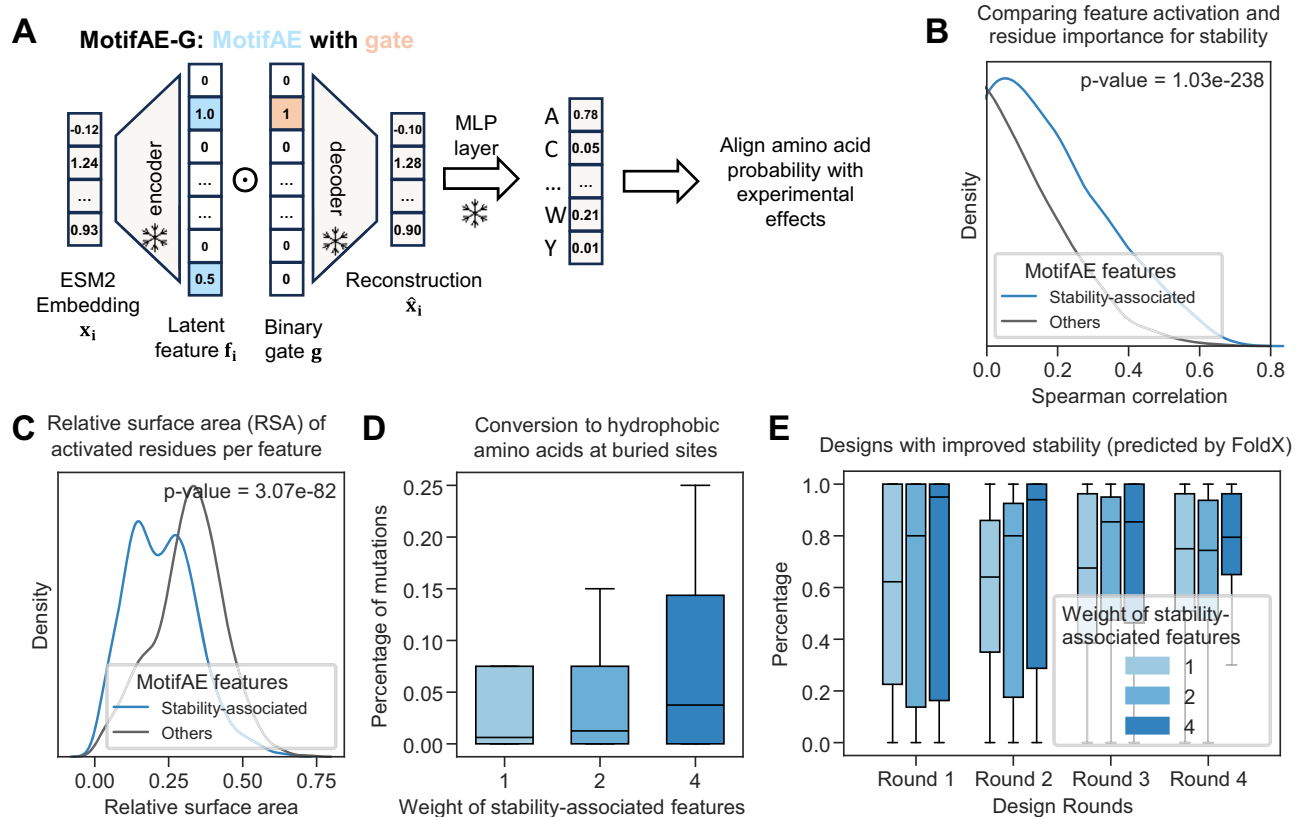
value =  $1.81\text{e} - 06$  ( $n = 65$ ); Organismal fitness:  $p$ -value =  $4.28\text{e} - 08$  ( $n = 62$ ); Activity:  $p$ -value =  $2.75\text{e} - 03$  ( $n = 25$ ); Expression:  $p$ -value =  $9.06\text{e} - 02$  ( $n = 15$ ); Binding:  $p$ -value =  $3.30\text{e} - 01$  ( $n = 10$ ). \*\*\*,  $p$ -value < 0.001; \*\*,  $p$ -value < 0.01; ns: not significant. **E** Spearman correlations of the best-match feature for residue importance in each DMS experiment. Each column shows the performance of MotifAE (dark dots) and SAE (light dots) for a single experiment. Dot size reflects the number of residues with at least ten assayed mutations, and colors indicate experimental categories. The five categories and the number of DMS experiments per category are listed below the plot. **F** Visualization of residue importance and MotifAE feature activations on protein structures. Upper panel: colors correspond to experimental categories, with darker colors indicating important residues with larger mean mutation effects. Lower panel: darker colors denote higher feature activation scores. Source data are provided as a Source Data file.

importance for stability, finding that their activation scores correlate more strongly with mean mutation effect on stability per residue than those of non-selected features (Fig. 6B). Next, we analyzed the relative solvent-accessible surface area and found that selected features preferentially activate at buried residues (Fig. 6C), which is consistent with the critical role of the domain core in maintaining structural stability. Lastly, we investigated whether steering selected features could enable stability-guided sequence redesign, using an iterative redesign sampling strategy<sup>29</sup> (see Methods). Analysis of the introduced mutations shows that increasing the weights of stability-associated features led to more substitutions toward hydrophobic residues at buried sites (Fig. 6D). Higher weights on these features also yield more redesigns

with improved stability relative to the original domains across redesign rounds, as predicted by FoldX<sup>28</sup> (Fig. 6E and S6C). Overall, these results demonstrate that aligning latent features with experimental data enables MotifAE-G to identify biophysically meaningful features associated with domain folding stability.

## Discussion

In this study, we present MotifAE, an unsupervised framework for extracting interpretable functional sequence patterns from pLMs. By incorporating a smoothness loss, MotifAE encourages coherent feature activation, which enhances its ability to discover functional motifs. The smoothness loss also propagates along the protein



**Fig. 6 | Predicting a stability-specific fitness landscape via supervised MotifAE feature selection.** **A** Schematic of applying MotifAE-G framework to DMS data. A learnable binary gate was used to select latent features to align predicted LLR with experimental mutation effect; during training, only the gate parameters were updated. **B** Distribution of Spearman correlations between feature activation scores and mean mutation effects per residue in each protein. All features were compared against all DMS experiments, and only features activated in at least 30% of residues within a protein were included. Only positive correlations are shown. Two-sided Welch's t-test:  $p$ -value =  $1.03\text{e-}238$  ( $n = 4887$  for stability-associated feature-DMS pairs and  $n = 7171$  for others). **C** Mean relative solvent accessibility (RSA) of activated residues for stability-associated versus other features, only features with at least 50 activated residues in 412 proteins were analyzed. Two-sided Welch's t-test:  $p$ -value =  $3.07\text{e-}82$  ( $n = 1313$  for stability-associated and  $n = 1591$  for others). **D,E** Designing proteins with improved stability using iterative redesign sampling.

Gate weights of 1, 2, and 4 were applied to stability-associated features (shown in progressively darker blue), while other features were fixed at a gate weight of 1. **D** Buried residues were defined as sites with RSA < 0.2. All mutations from four rounds of redesign were included in the analysis. One-sided paired Wilcoxon signed-rank test: weight 1 vs. 2,  $p$ -value =  $2.77\text{e-}01$ ; weight 1 vs. 4,  $p$ -value =  $4.28\text{e-}02$ ; weight 2 vs. 4,  $p$ -value =  $7.49\text{e-}03$ ;  $n = 16$  for all tests. **E** Percentage of designs showing improved stability compared to their primary (unmutated) domains. Increasing the weight of stability-associated features shows a positive trend toward improved stability across design rounds (one-sided mixed-effects model,  $\beta = 0.029$ ,  $p$ -value =  $0.055$ ,  $n = 16$ ). The box extends from the first quartile to the third quartile, with a line marking the median, whiskers span to the most extreme points within  $1.5\times$  the interquartile range. Source data are provided as a Source Data file.

sequence, enabling the capture of some longer sequence patterns, and MotifAE feature activation scores correlate with residue importance for different domain functions and properties. Furthermore, we introduce MotifAE-G, a supervised extension that aligns latent features with experimental data, enabling the identification of features associated with a specific function or property and the improvement of predictive performance on related tasks.

While some SAE models have been trained for pLMs<sup>16–18</sup>, they directly adopted strategies developed for large language models (LLMs). Our work introduces methodological advances in both training and interpretation. For training, we introduce an additional smoothness loss, defined as the L1 norm of latent feature differences between neighboring residues, sharing the same form as the sparsity L1 norm. Adding the smoothness loss preserves latent space sparsity, maintains reconstruction quality, and promotes coherent activation of latent features, facilitating the identification of functional sequence patterns. For interpretation, a common approach is to compare activated residues with known functional sites, but this is limited by incomplete annotation databases and cannot reveal novel motifs. Although Simon et al.<sup>17</sup> used LLMs to annotate latent features, their

approach still relies on known protein functions in literatures. We provide two strategies for annotating and potentially discovering novel motifs. First, we calculate PSSM for each MotifAE feature; features with well-defined PSSM sequence patterns may correspond to novel motifs. Future studies are needed to systematically link these sequence patterns to biological functions. Second, our MotifAE-G framework can be applied to annotate latent features using experimental datasets.

We applied MotifAE-G to the stability DMS dataset<sup>26</sup> because it was generated in a single laboratory and is large enough to enable reliable identification of features associated with folding stability. We note that the MotifAE-G framework can be applied more broadly. A general MotifAE-G application consists of two steps. First, a task-specific model is trained to predict experimental measurements from embeddings, with MotifAE reconstructions serve as a drop-in replacement for pLM embeddings; many existing models can be directly used (here, we use the ESM2 MLP layer). Second, a gate, which can be binary or continuous, selectively amplifies or attenuates MotifAE feature activations. This gate is trained on the same task to modulate both reconstructed embeddings and downstream predictions. Following

training, the gate weights reflect the feature association with the measured function or property, and the activated residues of associated features may underlie the experimental measurement. Applying MotifAE-G to other large-scale experiments offers the potential to identify features associated with other biological functions and properties.

For benchmarking, we primarily compared MotifAE against the standard SAE. We note that Simon et al.<sup>17</sup> already demonstrated that an SAE trained on pLM embeddings substantially outperforms both the raw pLM embeddings and an SAE trained on embeddings of pLM with shuffled weights in terms of interpretability. Therefore, we did not include those controls here, as MotifAE outperforms the standard SAE. For motif region identification, no current model can perform systematic, unsupervised discovery of motifs to our knowledge, so no alternative methods are available for comparison. For residue importance for domain function and property, conservation scores derived from either multiple sequence alignments or pLM predictions perform comparably to, or even better than, best-match single MotifAE latent feature. However, these conservation scores are not interpretable: when a residue is conserved, the underlying reason—whether due to stability, binding, post-translational modification, or noise—cannot be determined. In contrast, by selecting a subset of MotifAE features associated with specific properties, we can attribute the pLM-predicted conservation score into interpretable components, enabling better prediction of those properties.

Looking forward, MotifAE could be further improved in several directions. First, the current smoothness loss operates at the sequence level, incorporating three-dimensional proximity could enable the capture of complex structural domains and their functional organization. Second, we currently use the L1 norm for the sparsity loss because it shares the same form as the smoothness loss. Future works are needed to explore how alternative sparsity-promoting strategies, such as TopK<sup>30</sup> and BatchTopK<sup>31</sup>, can be integrated with the smoothness constraint. Third, since deep learning models inherently bias to their training data, both the pre-training dataset of the underlying pLM and the dataset used to train MotifAE influence the biological signals captured by latent features. Investigating how dataset composition shapes latent features is important. Fourth, MotifAE features exhibit different levels of granularity, understanding and controlling this granularity could improve feature interpretability and downstream utility. Lastly, our protein redesign results were evaluated *in silico*; experimental validation will be important in the future to assess MotifAE's potential for protein redesign.

Overall, we established a framework for unsupervised discovery and interpretation of functional sequence patterns from pLMs, enabling the analysis of sequence patterns at evolutionary scale. Our work also contributes to the broader goal of making biological deep learning models more interpretable. By systematically uncovering the sequence determinants of different protein functions and properties, MotifAE shows the potential to advance protein annotation, mutation effect interpretation, and rational protein engineering.

## Methods

### Sparse autoencoder architecture

Sparse autoencoder (SAE) projects language model embeddings into a sparse latent space. The model can be formulated as:

$$f = \text{ReLU}(W_{\text{encoder}}(x - b)) \quad (1)$$

$$\hat{x} = W_{\text{decoder}} \cdot f + b \quad (2)$$

where  $W_{\text{encoder}}$  projects the original embedding  $x$  into the high dimensional latent space  $f$  and  $W_{\text{decoder}}$  does the reverse to get the reconstruction  $\hat{x}$ .  $b$  is the bias term shared between encoder and decoder. ReLU was used as the activation function, setting negative

values to zero and retaining positive values. Both the encoder and decoder consist of a single linear layer.

Here, the 650-million-parameter version of ESM2 was used, which produces embeddings with a dimensionality of 1280. The SAE and MotifAE models were trained with varying latent space dimensions, and a dimension of 40,960 was selected (Fig. S1). We trained MotifAE models using both normalized (along embedding dimension) and raw ESM2 embeddings and found that models trained on raw embeddings performed better, particularly for fitness prediction. Therefore, all main figures are presented using the MotifAE model trained on raw embeddings. In Fig. S1, results using MotifAE models trained on normalized embeddings are presented.

### MotifAE loss function

MotifAE employs three losses: the reconstruction loss ( $L_{\text{rec}}$ ), the sparsity loss ( $L_{\text{sp}}$ ), and the smoothness loss ( $L_{\text{smooth}}$ ). The reconstruction and sparsity losses follow the SAE framework, preserving essential information from the input embeddings while enforcing sparse feature activations. The smoothness loss further encourages at least one nearby residue to have a similar latent feature activation, promoting coherent activations that capture the sequential continuity of protein motifs and basic structural elements. The loss functions are defined as follows:

$$L_{\text{rec}} = \frac{1}{n} \sum_{i=1}^n \|x_i - \hat{x}_i\|_2^2 \quad (3)$$

$$L_{\text{sp}} = \frac{1}{n} \sum_{i=1}^n |f_i| \quad (4)$$

$$D_i = \{|f_{i+1} - f_i|, \dots, |f_{i+k} - f_i|\} \quad (5)$$

$$L_{\text{smooth}} = \frac{1}{n-k} \sum_{i=1}^{n-k} \text{softmax}(D_i) \cdot D_i \quad (6)$$

$$L_{\text{total}} = L_{\text{rec}} + \alpha \times L_{\text{sp}} + \beta \times L_{\text{smooth}} \quad (7)$$

Here,  $i$  denotes the residue index;  $n$  denotes the protein length;  $x$  and  $\hat{x}$  represent the raw and reconstructed embedding vectors, respectively; and  $f$  denotes the latent feature activation vector.  $D_i$  quantifies the L1-norm difference in latent feature activations between residue  $i$  and its  $k$  neighboring residues.  $L_{\text{smooth}}$  encourages at least one nearby residue to have a similar latent feature activation. To achieve this, we apply a SoftMin over  $D_i$  to approximate the minimum difference of latent feature activation. Only one direction along the sequence is considered for each residue, but since  $L_{\text{smooth}}$  is computed for the full sequence, the opposite direction is accounted for when evaluating neighboring residues.  $k$  is the local window size used for similarity calculation.

### Model training

MotifAE and SAE were trained and evaluated on 2.3 million representative protein sequences clustered from the AlphaFold protein structure database using Foldseek<sup>27</sup> (downloaded from <https://afdb-cluster.steinerlab.workers.dev/>). Of these, 2.2 million sequences were used for training and the remainder for evaluation. Since ESM2 was trained with a maximum sequence length of 1,024, proteins exceeding this length were randomly truncated to a 1,024-residue region.

Models were implemented in PyTorch<sup>32</sup> (v2.4) and optimized using the Adam optimizer with a maximum learning rate of 0.001 and a 500-step linear warm-up. A batch size of 40 proteins was used. Residue order was preserved during training to retain positional information required by the smoothness loss. MotifAE models with different

window sizes were all trained with both sparsity and smoothness loss weights set to 0.04, whereas the standard SAE was trained with a sparsity loss weight of 0.085 and no smoothness loss. These loss weights were selected by testing multiple configurations to achieve a comparable number of activated latent features and similar reconstruction quality for both MotifAE and SAE. The sparsity and smoothness loss weights were gradually annealed over 5000 steps. Training was conducted for two epochs, and checkpoints at 80,000 steps for both MotifAE and SAE were used by considering latent sparsity and reconstruction quality. Each model required approximately 10 h of training on a single NVIDIA L40S GPU.

### Mutation effect prediction and data processing

The wild-type marginal method<sup>19</sup> was used to calculate log-likelihood ratio (LLR) to estimate mutation effects, which was calculated as:

$$LLR_i^{mt} = \log(p(x_i^{mt}, |, \mathbf{x})) - \log(p(x_i^{wt}, |, \mathbf{x})) \quad (8)$$

Where  $i$  represents residue index,  $mt$  and  $wt$  represent the mutant and wild-type amino acids,  $\mathbf{x}$  is the full sequence without mask.

The ProteinGYM<sup>20</sup> database comprises 216 deep mutational scanning (DMS) experiments spanning diverse functions and properties. In this study, only single substitution mutations were analyzed. For the domain functional organization analysis, mean mutation effects were computed for each residue, and only residues with at least ten assayed mutations were retained. Proteins containing fewer than twenty such residues were excluded, resulting in 180 DMS experiments. For each protein, only latent features that were activated in at least 30% of residues were analyzed. Three proteins lacking any SAE features meeting this threshold were excluded, yielding 177 proteins in the final analysis. The “DMS\_score” values were inverted for Spearman correlation analysis, as smaller values in ProteinGYM correspond to larger mutation effects.

The ClinVar<sup>23</sup> database provides expert-annotated mutations related to human diseases. ClinVar benign, likely benign, pathogenic, and likely pathogenic variants with at least one-star review status and no conflicting submissions were used<sup>33</sup>. For this study, we considered only genes containing at least five pathogenic / likely pathogenic and five benign / likely benign mutations, equal number of pathogenic / likely pathogenic and benign / likely benign mutations were randomly sampled from each protein, resulting in a gene-balanced dataset of 2272 pathogenic / likely pathogenic and 2272 benign / likely benign mutations in 207 genes. We used this dataset because ClinVar exhibits strong gene-level bias: the ratio of pathogenic to benign mutations varies substantially among genes<sup>34</sup>.

For protein stability<sup>26</sup>, the dataset “Tsuboyama2023\_Dataset2\_Dataset3\_20230416.csv” and AlphaFold2-predicted protein structures were downloaded from <https://zenodo.org/records/7992926>. Proteins were defined using the “WT\_name” column in the table. Only single substitution mutations with “ddG\_ML” values were analyzed. The “ddG\_ML” values for the same sequence were averaged. Solvent accessibility was computed from AlphaFold2-predicted structures using mdtraj<sup>35</sup>. The raw solvent-accessible surface areas were normalized by the maximum accessible area of each amino acid type, resulting in normalized values ranging from 0 to 1. For the domain functional organization analysis, the “ddG\_ML” values were inverted for Spearman correlation analysis (Fig. 6B).

### Evaluation on ELM motifs and CATH domains

ELM motifs were downloaded from the ELM<sup>1</sup> database in January 2025, and only instances annotated as true positive were analyzed. For proteins longer than 1,024 residues, a region of length 1,024 was selected with the motif region positioned at the center. IUPred3<sup>36</sup> was used to predict disordered regions: residues with prediction score >

0.5 were classified as disordered, whereas others were classified as ordered.

CATH<sup>7</sup> v4.3.0 domain boundaries and annotations were downloaded. Sequences of PDB structure in CATH were mapped to UniProt<sup>37</sup> Reviewed proteins, and the longest PDB structure was retained for each UniProt protein. Only UniProt proteins with lengths no longer than 1,024 residues were analyzed. CATH domains were required to lie fully within the mapped PDB–UniProt region, be at least 30 residues long, span no more than 70% of the full protein length, and be present in at least five different proteins. After filtering, 400 CATH domains across 5,076 proteins were included in our analysis.

Both tasks were formulated as a binary classification problem in which each residue was labeled as either belonging to a motif/domain or not. In MotifAE and SAE, each residue was represented by a 40,960-dimensional latent feature vector. For each latent feature, its activation value was used to predict motif/domain residues or not. To assess significance, we used a structured permutation test that preserved the activation structure of each feature. For each feature, non-zero activation values were shuffled across all activation blocks (single or contiguous activations), and the blocks were then relocated to random positions along the sequence. We generated 5000 permutations of the best-matching feature per motif or domain. As fewer than 5000 features were activated in any protein set we evaluated here (typically a few hundred per protein), 5000 permutations provide a sufficiently stringent null distribution.

### Activated peptides and PSSM analyses

Activated peptides for each MotifAE feature were identified across the 2.3 million representative proteins. Only activated peptides with lengths between 5 and 30 residues were retained for motif-related analyses. To calculate position-specific scoring matrices (PSSMs), up to top 1000 activated peptides ranked by mean activation scores were used for each feature. GibbsCluster 2.0<sup>25</sup>, with parameters -g 1 -l 5 -T -j 1 -l 1 -D 1, was employed to align the activated peptides and remove outliers for each feature. The motif length was set to five amino acids by default. We also evaluated motif lengths of four and six amino acids. For the six-residue motifs, we used the top 1000 activated peptides with lengths ranging from 6 to 30 residues. A trash cluster was used to remove sequences that do not align well with others. The aligned regions generated by GibbsCluster were subsequently used to construct PSSMs, which are calculated simply as the frequencies of 20 amino acids. PSSMs were visualized using the Logomaker<sup>38</sup> Python package.

To compare ELM regular expression (regex) with PSSMs, each PSSM was slid along the regex to evaluate all possible alignments with at least three residues of overlap. For each alignment, the PSSM probabilities of residues permitted by the regex were analyzed, and the highest average across all valid alignments for each PSSM-regex pair was reported in Fig. 4D.

### MotifAE-G architecture and training

The MotifAE-G introduces a gate to amplify or attenuate feature activations during embeddings reconstruction, which were subsequently used for downstream tasks. The gate was applied to each latent feature of MotifAE as:

$$\mathbf{f}' = \mathbf{f} \odot \text{gate} \quad (9)$$

where  $\odot$  denotes element-wise multiplication.

For protein stability prediction, a learnable binary gate was introduced while keeping ESM2 and MotifAE parameters fixed. Each gate value could only take 0 or 1, features with gates value of 1 were retained for reconstructing embeddings. The gate was binarized using the Straight-Through Estimator<sup>39</sup>. In the forward pass, each gate was

discretized as:

$$\text{gate} = \begin{cases} 1 & \sigma(g) > 0.5 \\ 0 & \sigma(g) \leq 0.5 \end{cases} \quad (10)$$

while in the backward pass, gradients were propagated directly through the sigmoid output  $\sigma(g)$ .

The reconstructed embeddings from the selected features were passed through the ESM2 MLP layer to predict probabilities of 20 amino acids, which were then converted into LLRs between the mutant and wild-type amino acids (Eq. 8). The resulting LLRs were compared with mutation effect on protein stability “ddG\_ML” to compute a soft Spearman correlation loss:

$$L = -\frac{\text{Cov}(r(\text{ddG}), r(\text{LLR}))}{\sqrt{\text{Var}(r(\text{ddG}))}\sqrt{\text{Var}(r(\text{LLR}))}} + \lambda |g| \quad (11)$$

where  $r(\bullet)$  denotes the differentiable rank function<sup>40</sup>. The first term maximizes the rank correlation between predicted and experimental mutation effects, while the second imposes an L1 regularization on the gate values to constrain the number of selected features. The weighting coefficient  $\lambda$  was set to 1 based on performance on the validation set.

For MotifAE-G training, the 412 domains<sup>26</sup> were clustered using Foldseek<sup>27</sup> easy-cluster with the parameters: --min-seq-id 0.3, yielding 116 clusters. The resulting domain clusters were then split into training, validation, and test sets at a 3:1:1 ratio, yielding 249 domains from 69 clusters for training, 78 domains from 23 clusters for validation, and 85 domains from 24 clusters for testing.

Model training was implemented in PyTorch using the Adam optimizer with default parameters. Each batch contained a single protein, and gradients were accumulated across all training proteins before updating the gate parameters, meaning the parameters were updated once per epoch. Models were trained for 100 epochs, and the checkpoint with the highest Spearman correlation on the validation set was selected for downstream analyses.

We further benchmarked other models (Fig. S6B). (1) A supervised model built by fine-tuning the ESM2 lm\_head (the module that further processes embeddings before predicting amino acid probabilities, consisting of a 1280×1280 dense layer and layer normalization). (2) MotifAE-G model with continuous gates instead of binary gates, with 1280×32 trainable parameters (MotifAE latent feature dimension). (3) FoldX<sup>28</sup> (version 20251231), a widely used biophysics-based predictor. The first two models were trained using the same dataset and training strategy to MotifAE-G. For FoldX, the predicted structures were first repaired using the RepairPDB command, and the BuildModel command was then applied to compute  $\Delta\Delta G$  values. The comparison between supervised models (Fig. S6B) highlights the importance of model capacity when sufficient training data are available.

### MotifAE-G stability-guided protein redesign and stability prediction

pLMs can be used for protein redesign by sampling from their predicted amino acid probabilities. However, directly generating sequences from pLMs does not necessarily yield proteins with the desired property. To address this, several approaches have been proposed to fine-tune pLMs to guide the design process toward specific functions or properties<sup>41,42</sup>. Here, we explored whether MotifAE-G can be used for property-specific protein design. Specifically, we steered the stability-associated features identified from the DMS dataset to design proteins with improved stability.

Because ESM2 is not suitable for de novo sequence generation, we adopted an iterative redesign sampling strategy<sup>29</sup>: at each iteration, a single mutation was sampled according to amino acid probabilities predicted by MotifAE-G, with higher gate weights amplifying the

influence of selected features, while other features kept a weight of one. We applied weights of 1, 2, and 4 to selected features, where a weight of 1 served as the baseline, equivalent to redesign using ESM2. We sampled one protein from each of the 24 test clusters and performed four rounds of iterative redesigns. We did not conduct more rounds because introducing a few new interactions is typically sufficient to stabilize the protein, whereas excessive mutations may alter the native fold. In each round, the probabilities of all single substitutions were predicted using MotifAE-G, with different weights applied to stability-associated features. The wild-type sequence was used without masking, all mutations were ranked by their relative probability compared to the wild type:  $p(x_i^{mt}, l, \mathbf{x})/p(x_i^{wt}, l, \mathbf{x})$ . A top- $k$  sampling strategy<sup>29</sup> (with  $k = 10$  in this study) was applied to select one mutation, where the sampling probability of each mutation was proportional to its relative probability. The entire procedure was repeated 20 times for each domain.

FoldX<sup>28</sup> (version 20251231) was used to predict protein stability from the ESMFold<sup>8</sup> predicted structures. Specifically, the “RepairPDB” command was first applied, followed by the “Stability” command to compute folding energy. The sign of the total energy was reversed to represent  $\Delta G$ . Among the 24 representative proteins from the test clusters, 16 proteins with ESMFold pLDDT >0.85 and FoldX-predicted  $\Delta G > 2$  kcal/mol were selected for subsequent redesign. Similarly, only redesigned sequences meeting the same thresholds were analyzed to ensure structural reliability and meaningful stability estimation.

### Reporting summary

Further information on research design is available in the Nature Portfolio Reporting Summary linked to this article.

### Data availability

All data used in this work are publicly available. 2.3 million representative proteins were downloaded from <https://afdb-cluster.steineggerlab.workers.dev/>. ProteinGYM DMS data were downloaded from <https://proteingym.org/download>. ELM motifs were downloaded from <http://elm.eu.org/downloads.html>. CATH domains were downloaded from <https://www.cathdb.info/wiki?id=data:index>. PDB-UniProt sequence matches were downloaded from [https://ftp.ebi.ac.uk/pub/databases/msd/sifts/flatfiles/csv/uniprot\\_segments\\_observed.csv.gz](https://ftp.ebi.ac.uk/pub/databases/msd/sifts/flatfiles/csv/uniprot_segments_observed.csv.gz). For protein stability, the DMS data and AlphaFold2-predicted protein structures were downloaded from <https://zenodo.org/records/7992926><sup>43</sup>. The data are available upon request. Source data are provided with this paper.

### Code availability

Codes are available at GitHub: <https://github.com/CHAOHOU-97/MotifAE.git>. Model weights can be downloaded from: <https://doi.org/10.5281/zenodo.20547228> (<https://zenodo.org/records/20547228>)<sup>44</sup>.

### References

- Kumar, M. et al. ELM-the Eukaryotic Linear Motif resource-2024 update. *Nucleic Acids Res.* **52**, D442–D455 (2024).
- Tokheim, C. et al. Systematic characterization of mutations altering protein degradation in human cancers. *Mol. Cell* **81**, 1292–1308 e11 (2021).
- Savojardo, C., Martelli, P. L. & Casadio, R. Finding functional motifs in protein sequences with deep learning and natural language models. *Curr. Opin. Struct. Biol.* **81**, 102641 (2023).
- Hou, C., Li, Y., Wang, M., Wu, H. & Li, T. Systematic prediction of degrons and E3 ubiquitin ligase binding via deep learning. *BMC Biol.* **20**, 162 (2022).
- Berman, H. M. et al. The protein data bank. *Nucleic Acids Res.* **28**, 235–242 (2000).

6. Lau, A. M. et al. Exploring structural diversity across the protein universe with the encyclopedia of domains. *Science* <https://doi.org/10.1126/science.adq4946> (2024).
7. Waman, V. P. et al. CATH 2024: CATH-alphaflow doubles the number of structures in CATH and Reveals Nearly 200 New Folds. *J. Mol. Biol.* **436**, 168551 (2024).
8. Lin, Z. et al. Evolutionary-scale prediction of atomic-level protein structure with a language model. *Science* **379**, 1123–1130 (2023).
9. Jumper, J. et al. Highly accurate protein structure prediction with AlphaFold. *Nature* <https://doi.org/10.1038/s41586-021-03819-2> (2021).
10. Kulmanov, M. et al. Protein function prediction as approximate semantic entailment. *Nat. Mach. Intell.* **6**, 220–228 (2024).
11. Zhang, Z. et al. Protein language models learn evolutionary statistics of interacting sequence motifs. *Proc. Natl. Acad. Sci. USA* **121**, e2406285121 (2024).
12. Vig, J. et al. Bertology meets biology: interpreting attention in protein language models. (2020).
13. Rives, A. et al. Biological structure and function emerge from scaling unsupervised learning to 250 million protein sequences. *Proc. Natl. Acad. Sci. USA* **118**, 239118 (2021).
14. Cunningham, H., Ewart, A., Riggs, L., Huben, R. & Sharkey, L. Sparse autoencoders find highly interpretable features in language models. Preprint at <https://doi.org/10.48550/arXiv.2309.08600> (2023).
15. Yun, Z., Chen, Y., Olshausen, B. A. & LeCun, Y. Transformer visualization via dictionary learning: contextualized embedding as a linear superposition of transformer factors. Preprint at <https://doi.org/10.48550/arXiv.2103.15949> (2023).
16. Adams, E., Bai, L., Lee, M., Yu, Y. & AlQuraishi, M. From mechanistic interpretability to mechanistic biology: training, evaluating, and interpreting sparse autoencoders on protein language models. 2025.02.06.636901 Preprint at <https://doi.org/10.1101/2025.02.06.636901> (2025).
17. Simon, E. & Zou, J. InterPLM: discovering interpretable features in protein language models via sparse autoencoders. *Nat. Methods* **22**, 2107–2117 (2025).
18. Gujral, O., Bafna, M., Alm, E. & Berger, B. Sparse autoencoders uncover biologically interpretable features in protein language model representations. *Proc. Natl. Acad. Sci.* **122**, e2506316122 (2025).
19. Brandes, N., Goldman, G., Wang, C. H., Ye, C. J. & Ntranos, V. Genome-wide prediction of disease variant effects with a deep protein language model. *Nat. Genet.* **55**, 1512–1522 (2023).
20. Notin, P. et al. ProteinGym: large-scale benchmarks for protein fitness prediction and design. *Adv. Neural Inf. Process. Syst.* **36**, 64331–64379 (2023).
21. Hou, C., Liu, D., Zafar, A. & Shen, Y. Understanding Language Model Scaling on Protein Fitness Prediction. 2025.04.25.650688 Preprint at <https://doi.org/10.1101/2025.04.25.650688> (2025).
22. Barrio-Hernandez, I. et al. Clustering predicted structures at the scale of the known protein universe. *Nature* **622**, 637–645 (2023).
23. Landrum, M. J. et al. ClinVar: updates to support classifications of both germline and somatic variants. *Nucleic Acids Res.* **53**, D1313–D1321 (2025).
24. Okada, M. Regulation of the Src family kinases by Csk. *Int. J. Biol. Sci.* **8**, 1385–1397 (2012).
25. Andreatta, M., Alvarez, B. & Nielsen, M. GibbsCluster: unsupervised clustering and alignment of peptide sequences. *Nucleic Acids Res.* **45**, W458–W463 (2017).
26. Tsuboyama, K. et al. Mega-scale experimental analysis of protein folding stability in biology and design. *Nature* **620**, 434–444 (2023).
27. van Kempen, M. et al. Fast and accurate protein structure search with Foldseek. *Nat. Biotechnol.* <https://doi.org/10.1038/s41587-023-01773-0> (2023).
28. Schymkowitz, J. et al. The FoldX web server: an online force field. *Nucleic Acids Res.* **33**, W382–W388 (2005).
29. Darmawan, J. T., Gal, Y. & Notin, P. Sampling protein language models for functional protein design. (2025).
30. Gao, L. et al. Scaling and evaluating sparse autoencoders. in (2024).
31. Bussmann, B., Leask, P. & Nanda, N. BatchTopK Sparse Autoencoders. Preprint at <https://doi.org/10.48550/arXiv.2412.06410> (2024).
32. Adam Paszke et al. PyTorch: an imperative style, high-performance deep learning library. in *Proceedings of the 33rd International Conference on Neural Information Processing Systems* Article 721 (Curran Associates Inc., 2019).
33. Zhong, G., Zhao, Y., Zhuang, D., Chung, W. K. & Shen, Y. PreMode predicts mode-of-action of missense variants by deep graph representation learning of protein sequence and structural context. *Nat. Commun.* **16**, 7189 (2025).
34. Cheng, J. et al. Accurate proteome-wide missense variant effect prediction with AlphaMissense. *Science* <https://doi.org/10.1126/science.adg7492> (2023).
35. McGibbon, R. T. et al. MDTraj: a modern open library for the analysis of molecular dynamics trajectories. *Biophys. J.* **109**, 1528–1532 (2015).
36. Erdos, G., Pajkos, M. & Dosztanyi, Z. IUPred3: prediction of protein disorder enhanced with unambiguous experimental annotation and visualization of evolutionary conservation. *Nucleic acids Res.* **49**, W297–W303 (2021).
37. UniProt, C. UniProt: a worldwide hub of protein knowledge. *Nucleic Acids Res.* **47**, D506–D515 (2019).
38. Tareen, A. & Kinney, J. B. Logomaker: beautiful sequence logos in Python. *Bioinformatics* **36**, 2272–2274 (2020).
39. Yin, P. et al. Understanding straight-through estimator in training activation quantized neural nets. Preprint at <https://doi.org/10.48550/arXiv.1903.05662> (2019).
40. Blondel, M., Teboul, O., Berthet, Q. & Djolonga, J. Fast differentiable sorting and ranking. Preprint at <https://doi.org/10.48550/arXiv.2002.08871> (2020).
41. Dieckhaus, H., Brocidiaco, M., Randolph, N. Z. & Kuhlman, B. Transfer learning to leverage larger datasets for improved prediction of protein stability changes. *Proc. Natl. Acad. Sci. USA* **121**, e2314853121 (2024).
42. Madani, A. et al. Large language models generate functional protein sequences across diverse families. *Nat. Biotechnol.* **41**, 1099–1106 (2023).
43. Tsuboyama, K. et al. Mega-scale experimental analysis of protein folding stability in biology and design. <https://doi.org/10.5281/zenodo.7992926> (2023).
44. Hou, C. Data and model weights for MotifAE, a sparse autoencoder for unsupervised discovery of functional motifs from protein language model. <https://doi.org/10.5281/zenodo.20547228> (2026).

## Acknowledgements

We thank Aziz Zafar and other members of the Shen Lab for helpful discussions.

## Author contributions

Y.S. and C.H. conceived the study. C.H. designed and implemented the experiments. D.L. assisted with mutation analysis. All authors evaluated and interpreted the results and approved the final version of the manuscript.

## Funding statement

Y.S. discloses support from NIH grants R35GM149527 and Simons Foundation SFARI #1019623. The funders had no role in study design, data collection and analysis, decision to publish or preparation of the manuscript. The other authors declare no relevant funding.

## Competing interests

The authors declare no competing interests.

## Additional information

**Supplementary information** The online version contains supplementary material available at <https://doi.org/10.1038/s41467-026-77333-2>.

**Correspondence** and requests for materials should be addressed to Chao Hou or Yufeng Shen.

**Peer review information** *Nature Communications* thanks Alessio Ansuini, Haiyuan Yu and other anonymous reviewer for their contribution to the peer review of this work. A peer review file is available.

**Reprints and permissions information** is available at <http://www.nature.com/reprints>

**Publisher's note** Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

**Open Access** This article is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License, which permits any non-commercial use, sharing, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if you modified the licensed material. You do not have permission under this licence to share adapted material derived from this article or parts of it. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by-nc-nd/4.0/>.

© The Author(s) 2026

This Article in Press is shared early to give you faster access to new research. It is citable and carries a permanent DOI. The final edited version will replace it automatically.