

mAbs

ISSN: 1942-0862 (Print) 1942-0870 (Online) Journal homepage: www.tandfonline.com/journals/kmab20

Prediction of protein biophysical traits from limited data: a case study on nanobody thermostability through NanoMelt

Aubin Ramon, Mingyang Ni, Olga Predeina, Rebecca Gaffey, Patrick Kunz, Shimobi Onuoha & Pietro Sormanni

To cite this article: Aubin Ramon, Mingyang Ni, Olga Predeina, Rebecca Gaffey, Patrick Kunz, Shimobi Onuoha & Pietro Sormanni (2025) Prediction of protein biophysical traits from limited data: a case study on nanobody thermostability through NanoMelt, mAbs, 17:1, 2442750, DOI: 10.1080/19420862.2024.2442750

To link to this article: <https://doi.org/10.1080/19420862.2024.2442750>



© 2025 The Author(s). Published with license by Taylor & Francis Group, LLC.



[View supplementary material](#)



Published online: 08 Jan 2025.



[Submit your article to this journal](#)



Article views: 5374



[View related articles](#)



[View Crossmark data](#)



Citing articles: 3 [View citing articles](#)

REPORT



Prediction of protein biophysical traits from limited data: a case study on nanobody thermostability through NanoMelt

Aubin Ramon ^a, Mingyang Ni^a, Olga Predeina^a, Rebecca Gaffey^a, Patrick Kunz^{b*}, Shimobi Onuoha^c, and Pietro Sormanni ^a

^aCentre for Misfolding Diseases, Yusuf Hamied Department of Chemistry, University of Cambridge, Cambridge, UK; ^bDivision of Functional Genome Analysis, German Cancer Research Center (DKFZ), Heidelberg, Germany; ^cChimeris UK, The Works, Cambridge, UK

ABSTRACT

In-silico prediction of protein biophysical traits is often hindered by the limited availability of experimental data and their heterogeneity. Training on limited data can lead to overfitting and poor generalizability to sequences distant from those in the training set. Additionally, inadequate use of scarce and disparate data can introduce biases during evaluation, leading to unreliable model performances being reported. Here, we present a comprehensive study exploring various approaches for protein fitness prediction from limited data, leveraging pre-trained embeddings, repeated stratified nested cross-validation, and ensemble learning to ensure an unbiased assessment of the performances. We applied our framework to introduce NanoMelt, a predictor of nanobody thermostability trained with a dataset of 640 measurements of apparent melting temperature, obtained by integrating data from the literature with 129 new measurements from this study. We find that an ensemble model stacking multiple regression using diverse sequence embeddings achieves state-of-the-art accuracy in predicting nanobody thermostability. We further demonstrate NanoMelt's potential to streamline nanobody development by guiding the selection of highly stable nanobodies. We make the curated dataset of nanobody thermostability freely available and NanoMelt accessible as a downloadable software and webserver.

ARTICLE HISTORY

Received 4 October 2024
Revised 10 December 2024
Accepted 11 December 2024

KEYWORDS

Protein fitness; thermostability; antibody engineering; antibody design; nanobody; machine learning; semi-supervised learning; ensemble model

CLASSIFICATION

Biological sciences – biophysics and computational biology

Introduction

Protein fitness, defined as the ability of a protein to perform its biological function, encompasses various functional and biophysical properties, including binding affinity, enzymatic activity, solubility, and conformational stability. In silico prediction and optimization of protein fitness have become increasingly prevalent, as they reduce the need for time- and resource-intensive experiments.^{1,2} Coarse but rather reliable predictions of protein fitness can now be made ab initio using property-agnostic models. For instance, large language models like ESM^{3,4} can assign probabilities to whether a polypeptide sequence is protein-like and capable of folding and functioning, while models like AlphaFold⁵ can predict native protein states, potentially inferring function.

However, the quantitative prediction of biophysical traits that underpin fitness, such as enzymatic activity, binding affinity, or native-state stability, requires property-specific approaches. Energy-based or heuristics ab initio calculations remain challenging, as they need high-resolution structural information as input, and rely on computationally demanding calculations, or on rather inaccurate approximations, which hamper many applications.^{6–8} Therefore, accurate, rapid, practically useful predictions with the potential to replace experiments generally require rather large amounts of diverse

protein-specific or family-specific experimental data for the training of data-driven approaches.

A key limitation in this field is the scarcity of diverse experimental data, which hinders model generalizability to novel or distant sequences.⁹ Moreover, the available data are often heterogeneous and poorly consistent, as measurements are frequently performed using different techniques or under varying conditions. When such inconsistent data are aggregated, it further complicates the training of effective predictors, as the variability introduced by differing methodologies can obscure true biophysical signals and reduce model accuracy. This necessitates robust training and performance assessment procedures to optimize data use and prevent biases, which are crucial when dealing with limited data.^{10,11}

Various strategies have been used to improve predictions in the context of limited data.¹² Transfer learning, for instance, leverages protein representations from pre-trained models on large databases for supervised downstream tasks with fewer data.^{13,14} Few-shot learning and semi-supervised learning further refine these models with minimal labeled data.¹⁵ Ensemble learning, which combines multiple models, enhances robustness and accuracy.¹⁶ Techniques such as bagging and stacking reduce overfitting and improve generalization by integrating the strengths of individual models.¹⁷

CONTACT Pietro Sormanni  ps589@cam.ac.uk  Centre for Misfolding Diseases, Yusuf Hamied Department of Chemistry, University of Cambridge, Lensfield road, Cambridge CB2 1EW, UK

*Current address: Peptone Switzerland AG, Via Vincenzo Vela 5, 6500 Bellinzona, Switzerland

 Supplemental data for this article can be accessed online at <https://doi.org/10.1080/19420862.2024.2442750>

© 2025 The Author(s). Published with license by Taylor & Francis Group, LLC.

This is an Open Access article distributed under the terms of the Creative Commons Attribution License (<http://creativecommons.org/licenses/by/4.0/>), which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited. The terms on which this article has been published allow the posting of the Accepted Manuscript in a repository by the author(s) or with their consent.

In this work, we introduce a novel framework for learning protein biophysical traits from limited and heterogeneous labeled data. We first estimate the data requirements for effective generalization using unconstrained artificial data, and then apply this framework to develop NanoMelt, a predictor of nanobody thermostability. Thermostability, a crucial fitness attribute, quantifies a protein's resistance to denaturation upon heating and correlates strongly with conformational stability¹⁸ and resistance to physicochemical stresses.¹⁹ It significantly influences protein developability, affecting expression levels,²⁰ colloidal stability, aggregation propensity,²¹ and pharmacokinetics.²² Typically, quantified by the melting temperature (T_m),²³ thermostability is often determined via techniques such as differential scanning fluorimetry (DSF), differential scanning calorimetry (DSC), and circular dichroism (CD). Despite the existence of large datasets like ProThermDB,²⁴ which contains approximately 31,000 data points, models developed^{6,25–28} from these datasets often exhibit poor predictive performance for protein subclasses with limited representation, such as monoclonal antibodies (mAbs). Fine-tuning on antibody-specific measurements has been necessary to achieve satisfactory predictions in these cases.²⁹

Nanobodies, or single-domain antibodies, are small antibody fragments that typically exhibit good expressibility, solubility, stability, and a binding affinity comparable to full-length antibodies.³⁰ Since their discovery,³¹ nanobodies have found numerous biological applications^{32,33} and garnered substantial therapeutic interest,³⁴ especially following the approval of Caplacizumab and the rise of multi-specific and multi-paratopic drugs. However, public-domain data on nanobody thermostability are limited. The NbThermo³⁵ database, which contains 548 datapoints, is the most recent compilation of

nanobody melting temperatures. The heterogeneity in experimental conditions and methodologies across these measurements results in poorly consistent data, complicating the prediction of nanobody thermostability.

In this study, we present a comprehensive exploration of protein fitness learning with limited data, applying our framework to predict nanobody thermostability. We curated a novel dataset of 640 melting temperatures of unique nanobody sequences, by further cleaning and removing redundancies from NbThermo, merging it with other sources, and performing 129 new measurements. We developed NanoMelt by using an ensemble learning strategy that combines various regression techniques and sequence embeddings. Our approach ensures optimal data utilization and robust performance via a repeated stratified nested cross-validation procedure. We further demonstrate NanoMelt's potential to streamline nanobody development by guiding the selection and design of highly stable nanobodies. We have made the curated nanobody thermostability dataset freely available, and NanoMelt accessible both as downloadable software and user-friendly webserver.

Results

Preliminary search of minimal fitness dataset size

In this study, we introduce an extensive protein fitness prediction pipeline optimized for the use of limited data during training and performance assessment (Figure 1), which we apply to the prediction of nanobody thermostability. Training a regression model to accurately predict biophysical traits solely from sequence information remains challenging,

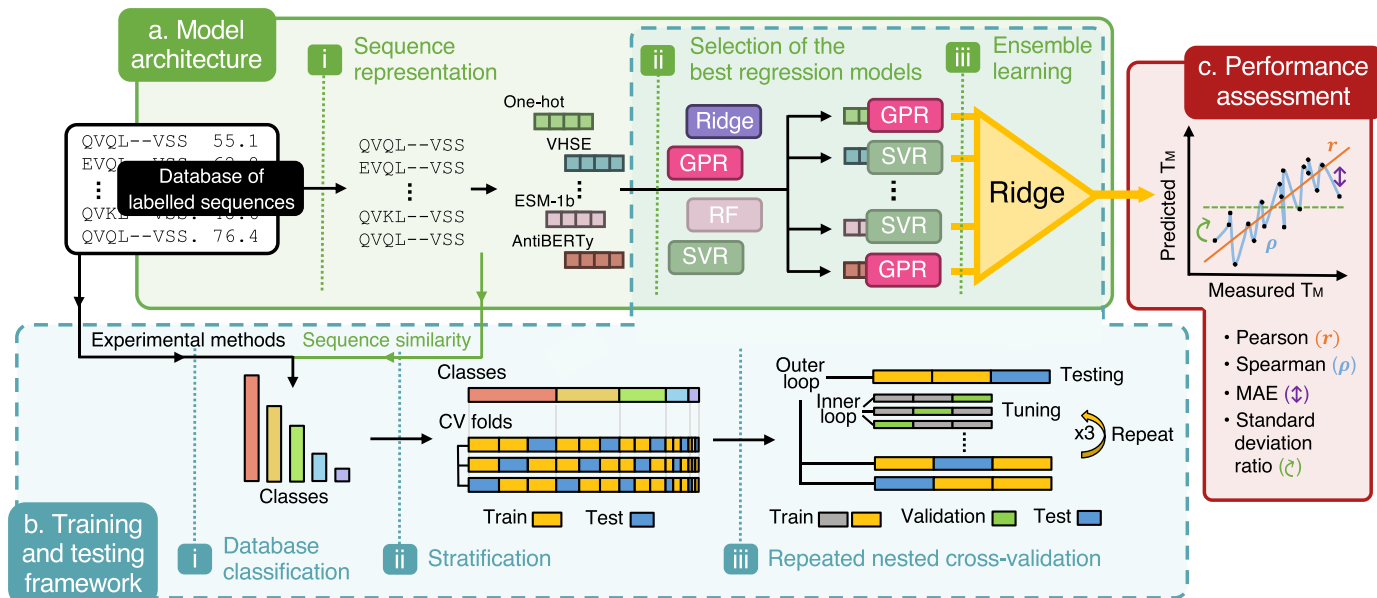


Figure 1. Framework for protein biophysical trait prediction with limited data. **(a)** Model architecture (green): sequences from a labeled protein database (black box) are represented using multiple embeddings (e.g., ESM-1b, one-hot, VHSE) and processed through various regression models (e.g., ridge, GPR, RF, SVR). The top-performing models for each embedding are combined into a ridge-based ensemble model. **(b)** Training and testing framework (blue): the dataset is clustered by experimental method (e.g., nanoDSF, DSF, DSC, CD in the example of thermostability measurements) and sequence similarity via k-medoids clustering. Stratification ensures consistent class distribution across training (yellow) and testing (blue) splits. Repeated nested cross-validation includes an inner loop for hyperparameter tuning (grey and green) within an outer loop for model testing (yellow and blue), repeated with three random seeds to report averaged performances and their standard deviations. **(c)** Performance assessment (red): Model performance is evaluated using Pearson's correlation (r), Spearman's correlation (p), mean absolute error (MAE), and standard deviation ratio (SDR), which indicates the model's tendency to regress towards the dataset mean – a common issue with limited data.

with performance largely influenced by the number of labeled sequences available and by how such sequences are represented numerically.

We first aimed to determine the minimum amount of labeled data necessary to train a reliable model capable of generalizing to sequences distant from those in the training set, and to evaluate how performance varies with different sequence representations or embeddings. Given the lack of an extensive dataset of nanobody thermostability measurements, we constructed an artificially labeled dataset. The advantage of using artificial data is the ability to generate it in a fully unconstrained manner, enabling the exploration of both large and small datasets of tailored diversity. A total of 10,000 nanobody sequences were selected, and their relative solubility was calculated using the CamSol algorithm³⁶ (see Materials and Methods). CamSol performs complex calculations, combining the physicochemical properties of amino acids to account for both short-range and long-range interactions, ultimately yielding a single solubility score for the whole sequence. We generated our artificial labeled data using two scores: the intrinsic solubility score, derived solely from the amino acid sequence, and the structurally corrected surface score, which also accounts for residue solvent exposure and proximity in the three-dimensional space, as obtained from modeled structures (see Materials and Methods). Due to the intricate and non-linear nature of these calculations, we hypothesized that these artificial datasets would approximate reasonably well real datasets of experimentally measured physicochemical properties, at least in terms of prediction complexity.

We trained ridge regression models via repeated cross-validation on datasets of varying sizes, ranging from 20 to 8,000 sequences, and tested all models on the same held-out test set of 2,000 sequences (see Materials and Methods, Fig. S1). Sequences were represented with eight different embeddings, including categorical (one-hot) and physicochemistry-based (VHSE³⁷) encodings of nanobody sequences aligned with the AHo numbering scheme to a fixed length of 149 positions, as well as pre-trained embeddings derived from protein (ESM-1b³⁸ and ESM-2⁴), antibody-specific (AbLang³⁹ and AntiBERTy⁴⁰), and nanobody-specific (nanoBERT⁴¹) large language models, and from a state-of-the-art nanobody structure predictor (NanobodyBuilder2⁴²), all of which were given non-aligned nanobody sequences as input (see Materials and Methods, Table S1).

As anticipated, increasing the training set size improved the Spearman's coefficient (ρ) on the held-out test set and reduced variance across repeated random splits (Fig. S1). The categorical one-hot and VHSE encodings exhibited greater overfitting, assessed by the discrepancy between performance on the training and test sets, and showed poorer ability to generalize, requiring a larger number of training sequences to reach a performance plateau on the test set. In contrast, pre-trained embeddings from large language models demonstrate superior generalization at lower training set sizes. At a training set size exceeding 600, the ESM-1b and ESM-2 begin to reach a performance plateau. For example, at a size of 640, the ESM-1b and ESM-2 show ρ values of 0.871 and 0.888, respectively, compared to 0.605 for the one-hot encoding on the

intrinsic solubility score (Fig. S1a). Meanwhile, models trained on AntiBERTy, AbLang, and nanoBERT embeddings achieved ρ values of 0.751, 0.793, and 0.802, respectively, suggesting that embeddings trained specifically on antibody or nanobody sequences do not perform as well as those derived from a generic, yet much larger, protein language model. The NanobodyBuilder2 embedding failed to provide sufficient information to the regression model, plateauing at a ρ value of only 0.352 with a training set size of 8,000 sequences. This poor performance is likely due to the NanobodyBuilder2 embedding's relatively low dimensionality, comprising 128 input features compared to 1,280 for ESM-1b and over 3,000 for one-hot encoding (149×21), as well as its training on a few thousand structures, in contrast to ESM-1b's 30 million sequences.

In conclusion, consistent with previous reports,^{14,16,43} semi-supervised learning using pre-trained embeddings is the most effective strategy for learning from small protein fitness datasets. Importantly, a training set with more than 600 labeled nanobody sequences appears to provide sufficient information for a sequence-based regression model to yield accurate and generalizable predictions of nanobody biophysical traits. Beyond this threshold, further increases in training data size yield only modest improvements in test-set performance.

Construction of the nanobody thermostability dataset

Apparent melting temperatures (T_m) of nanobodies have been sporadically reported in the literature. The NbThermo³⁵ database attempted to compile these thermostability datapoints from various publications. First, we added to the NbThermo a dataset with data from more recent mutational studies, obtaining 511 unique sequence datapoints (see Materials and Methods, Table S2).

However, our previous analysis on artificial data indicated that these may be insufficient to train a robust model for nanobody fitness prediction. To address this, we generated 129 additional datapoints by characterizing nanobodies specifically for this study (Figure 2). These nanobodies included both some already available in our laboratory, and a selection of diverse sequences not previously represented in the database. The largest contribution within NbThermo, comes from Kunz et al.,⁴⁴ who provided data for 64 diverse nanobodies, all measured under the same conditions, and for which we have the raw fluorescence readouts, meaning that we can re-fit these traces. Therefore, all new measurements were conducted under the same conditions as those used by Kunz et al.,⁴⁴ and all data were then analyzed together, thus yielding a self-consistent dataset of 193 datapoints with uniform technique, protein concentration, buffer, and pH (see Materials and Methods). As a control, we reproduced and characterized six nanobodies from Kunz et al., covering a broad range of T_m values. The T_m values from our measurements (Fig. S2) closely align with those from Kunz et al.,⁴⁴ confirming the high compatibility between the datasets. Measurements were performed at low concentrations to minimize the impact of heat-induced aggregation on equilibrium, which can cause shifts in the measured T_m ⁴⁴ (see

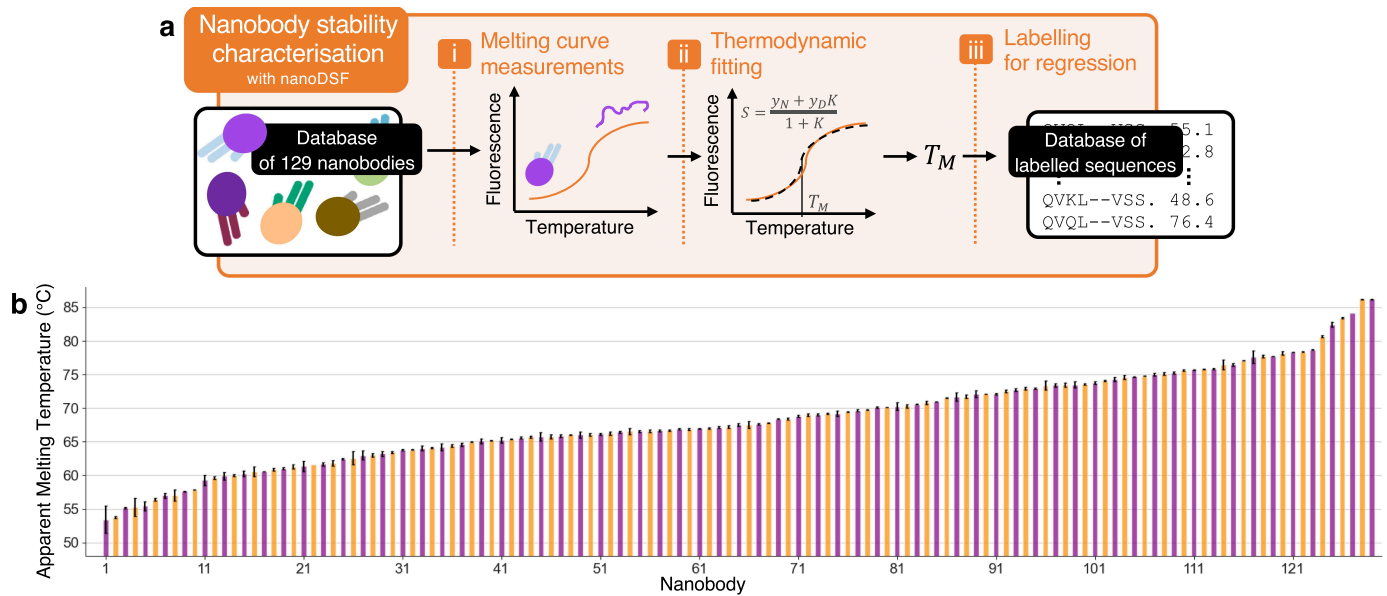


Figure 2. Expansion of nanobody thermostability data. **(a)** Nanobody thermostability characterization pipeline: the intrinsic fluorescence of 129 nanobodies was measured at increasing temperatures using nanoDSF. The resulting melting curves were fitted to a two-state protein denaturation model to determine the apparent melting temperatures (T_m) (see materials and methods). These T_m values are then used as labels to train and test our stability predictor, alongside measurements from the literature. **(b)** Fitted T_m for the 129 characterized nanobodies: bar heights are the mean T_m from duplicate measurements from separate runs, and error bars indicate the standard deviation. The average error is 0.26°C, indicating excellent consistency between measurements.

Materials and Methods, Figure 2a). We carefully extracted the melting temperatures by individually fitting fluorescence emission signals at 350 nm and 330 nm whenever feasible (Fig. S3). While fitting the 350/330 nm fluorescence ratio is common practice, this approach can introduce biases under certain conditions,⁴⁵ as illustrated here by a 3°C deviation observed in our analysis of the 7KBI melting curves (Fig. S4).

These additional measurements resulted in a final thermostability dataset comprising 640 distinct nanobodies, with T_m values ranging from 26.6°C to 98.2°C (Figure 3a). The small subset with $T_m \leq 42^\circ\text{C}$ comes from a single study of destabilizing mutations (see Materials and Methods). The dataset includes measurements from three experimental techniques: 251 from nanoDSF, 205 from CD, and 165 from sypro-orange DSF (Figure 3b).

We clustered the sequences using the k-medoids algorithm into 13 distinct groups (Figure 3c, Materials and Methods). A two-dimensional reduction of these clusters mapped onto the native nanobody space demonstrates the diversity of our dataset, with datapoints roughly evenly distributed across the sequence space (Figure 3d). Furthermore, the 640 sequences cover a similar range of CDR3 lengths (Fig. S5a) but exhibit a higher average somatic mutation rate compared to the sequences in the native dataset, likely due to the presence of engineered sequences in NanoMelt's dataset (Fig. S5b-c). Ultimately, the mutation distance matrix (see Materials and Methods, Fig. S6) reveals substantial sequence diversity within the dataset, with an average sequence identity of 64%, which is rather low for closely related proteins such as nanobodies.

Taken together, these results confirm that we have constructed a broad and diverse dataset of nanobody melting temperatures, achieved also by leveraging significant experimental efforts dedicated to characterizing 129 additional nanobodies under consistent conditions to maximize data quality. The complete dataset is available as Supplementary Dataset 1.

Selection of regression models

Our analysis using artificial data indicated that 640 datapoints should suffice for training robust models for nanobody fitness prediction. To assess the ability of different regression models to predict nanobody thermostability, we evaluated seven models: three linear (Ridge,⁴⁶ Huber,⁴⁷ and Elastic Net⁴⁸) and four non-linear (Random Forest⁴⁹ and LightGBM,⁵⁰ Support Vector Machine,⁵¹ and Gaussian Process⁵²). Each model used one of eight sequence embeddings as input (Table S1) to predict T_m .

The dataset is limited not only in size but also in consistency, with experimental and analytical variability contributing to differences in apparent-melting temperatures. For example, 82 sequences have multiple T_m measurements under varying conditions, resulting in a mean pairwise absolute difference of 2.4°C (Fig. S7). To mitigate biases from these inconsistencies and the limited data, we used a repeated stratified nested cross-validation procedure¹¹ (Figure 1).

Stratification maintained class distribution across folds, improving generalizability and reliability when training on imbalanced heterogeneous data.⁵³ We stratified based on the experimental method and sequence similarity cluster. Comparison between stratified and random cross-validation showed that stratification improved performance slightly but consistently and reduced variability (Fig. S8). Since diverse

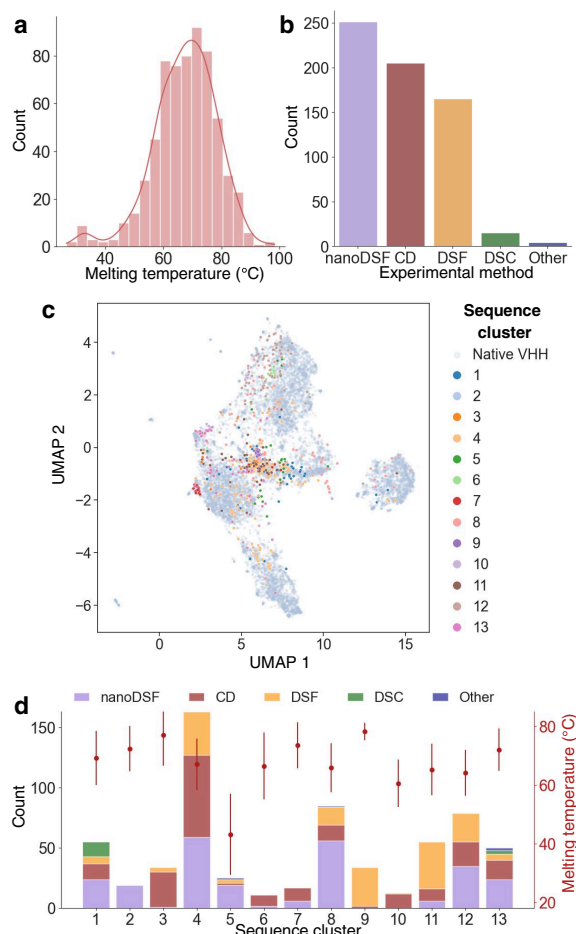


Figure 3. Overview of the nanobody thermostability dataset. **(a)** Histogram showing the distribution of apparent melting temperatures (°C) across the dataset. **(b)** Bar plot depicting the distribution of the experimental methods used. **(c)** UMAP projection ($n_neighbors = 20$) of the ESM-1b representation of 10,000 native nanobody sequences from the AbNatiV dataset (in grey) overlaid with the 640 sequences from our dataset. Sequences are colour-coded according to their k-medoids cluster number (see materials and methods). **(d)** Bar plot showing the distribution of experimental methods used within the 13 k-medoids sequence clusters (as coloured in panel c). The average melting temperature and standard deviation for each cluster are indicated by the red points (right y-axis). Further details on dataset's CDR3-length and germline diversity are provided in figure S5.

experimental methods were used to measure the thermostability in our database, stratification allows the distribution of these methods in the training data to be maintained. This ensures that the model learns from a more representative sample, reducing the risk of fitting to the most common experimental methods and compromising the ability to make predictions on the others.

Nested cross-validation allows the efficient use of limited data, while preventing data leakage between training, validation, and test sets, thus keeping the model fine-tuning independent from the model assessment.¹⁰ The procedure involves an inner loop for hyperparameter tuning and an outer loop for performance evaluation on an unseen test set, repeated with three random seeds to capture the magnitude of performance variations (Figure 1b).

Results across models and embeddings (Figure 4a, Tables 1 and S3) showed that Gaussian Process Regression (GPR) and

Support Vector Regression (SVR) consistently outperformed other models, with higher correlation and lower variance, though with a stronger tendency to overfit (e.g., SVR on ESM-1b: $\rho = 0.993 \pm 0.004$ for training vs. 0.816 ± 0.027 for testing). However, SVR and GPR are expected to fit closely the training data due to their use of kernel functions that model complex data relationships which can capture both the relevant patterns (signal) and the random variations (noise) within the dataset.

Among embeddings, the ESM-1b model demonstrated the best overall performance ($\rho = 0.816 \pm 0.027$ with SVR), surpassing those trained on antibody (AntiBERTy: $\rho = 0.745 \pm 0.036$ with SVR) and nanobody (nanoBERT: $\rho = 0.723 \pm 0.035$ with GPR) sequences. VHSE ($\rho = 0.797 \pm 0.027$ with SVR) and one-hot encoding ($\rho = 0.784 \pm 0.029$ with GPR) closely followed but were more prone to overfitting and to regression to the mean (as indicated by lower SDR, Table 1). This suggests that pre-trained embeddings offer superior generalization to sequences distant from the training set.

In summary, repeated stratified nested cross-validation offers a robust framework for comparing models trained on limited, heterogeneous data. Non-linear models with pre-trained embeddings achieved Spearman's correlations exceeding 0.8 on the test sets.

For each embedding, the table lists the best-performing regression model, ranked according to Spearman's coefficient. The Pearson's and Spearman's coefficients, MAE, and SDR are averaged across those on the test folds of the outer loop in the repeated nested cross-validation, and the reported uncertainties are the corresponding standard deviations. Ranks are provided across the 56 possible combinations of regression models and embeddings (see Table S3). The best performance in each column is highlighted in bold.

Ensemble learning

Prediction performance varied significantly depending on the combination of embeddings and regression models used. Ensemble learning, which combines multiple individually trained models into a single predictor, is known to enhance performance by leveraging model diversity, thereby increasing generalizability and robustness when working with limited, heterogeneous data.^{12,16}

In this study, predicted test-set T_m from individually trained regression models were used as inputs to train the ensemble model. We evaluated different combinations of models for ensemble learning with the same stratified nested cross-validation pipeline (see Materials and Methods, Figure 4c and Table S4). In addition to the top-ranking combinations (Table S3), we tested combinations of the best model for each different embedding (Table 1), and of different regression models trained on their best-performing embeddings (Table S3). We explored two ensemble strategies for these combinations: averaging the predictions from up to eight models (Table S3), and stacking them using a ridge regression meta-model⁵⁴ (see Materials and Methods). Both approaches

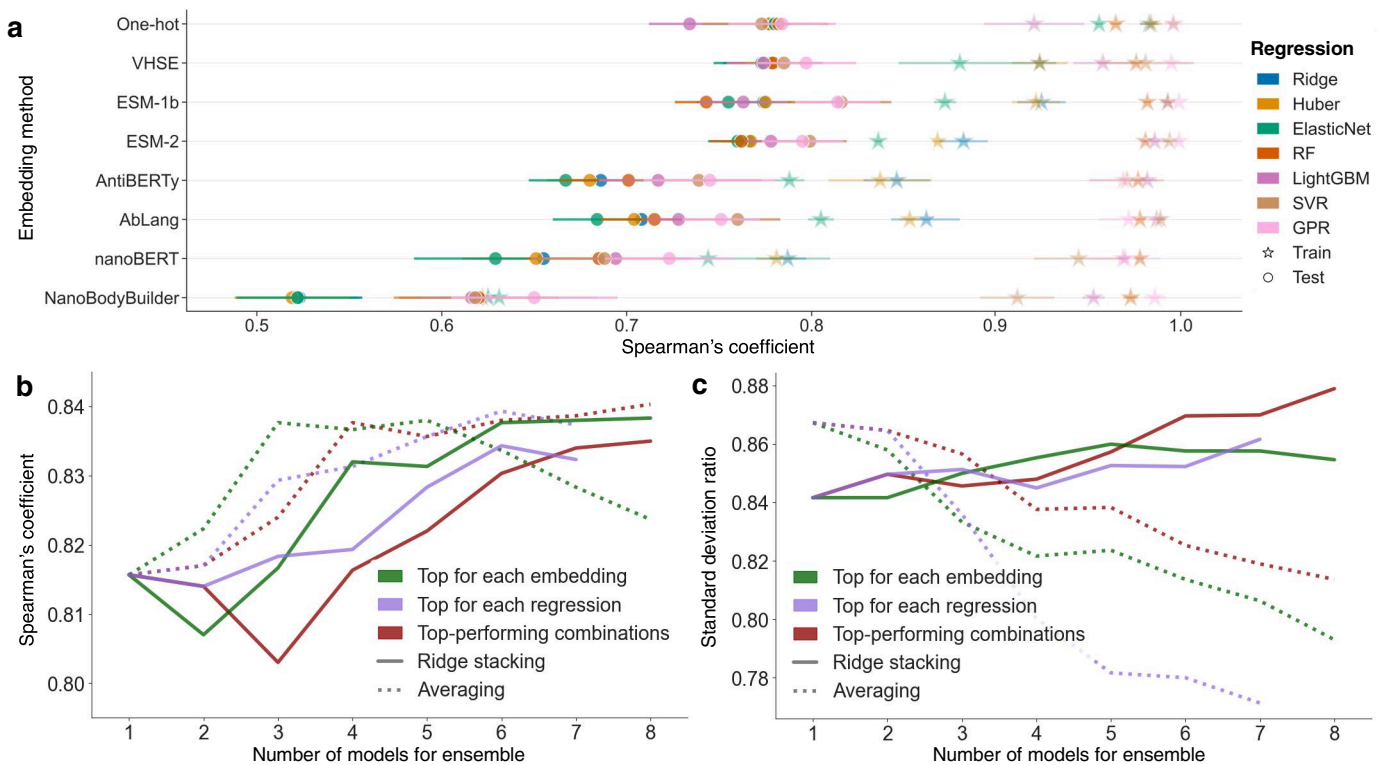


Figure 4. Performance evaluation of regression models for ensemble learning. **(a)** Spearman's coefficient for selected regression models across different embeddings using repeated nested stratified cross-validation. Each colour represents a regression method (see legend). Triangle markers indicate performance on the training set, while round markers indicate performance on the test set. Error bars denote standard deviations across cross-validation folds and repeats. **(b)** Test-set Spearman's coefficient and **(c)** SDR performances for the ridge stacking ensemble (in solid line) and the averaged ensemble (in dotted line) using input predictions from different model selections: top-performing models across diversified embeddings (green, Table 1); top-performing embeddings with diversified regression models (purple, table S3); and top-performing combinations of embedding and model from nested stratified cross-validation (red, table S3).

Table 1. Top-performing regression models for each sequence representation.

Rank (out of 56)	Embedding	Regression	Pearson's coefficient (r)	Spearman's coefficient (ρ)	MAE [$^{\circ}\text{C}$]	Standard deviation ratio (SDR)
1	ESM-1b	SVR	0.840 ± 0.028	0.816 ± 0.027	4.2 ± 0.3	0.87 ± 0.05
3	ESM-2	SVR	0.836 ± 0.025	0.799 ± 0.020	4.3 ± 0.3	0.87 ± 0.05
4	VHSE	GPR	0.823 ± 0.028	0.797 ± 0.027	4.5 ± 0.3	0.84 ± 0.04
7	One-hot	GPR	0.818 ± 0.031	0.784 ± 0.029	4.5 ± 0.3	0.82 ± 0.05
26	AbLang	SVR	0.800 ± 0.026	0.760 ± 0.023	4.8 ± 0.2	0.89 ± 0.05
29	AntiBERTy	GPR	0.786 ± 0.027	0.745 ± 0.036	5.0 ± 0.3	0.82 ± 0.06
34	nanoBERT	GPR	0.784 ± 0.027	0.723 ± 0.035	5.1 ± 0.4	0.81 ± 0.06
49	NanobodyBuilder2	GPR	0.744 ± 0.038	0.650 ± 0.045	5.4 ± 0.3	0.77 ± 0.04

For each embedding, the table lists the best-performing regression model, ranked according to Spearman's coefficient. The Pearson's and Spearman's coefficients, MAE, and SDR are averaged across those on the test folds of the outer loop in the repeated nested cross-validation, and the reported uncertainties are the corresponding standard deviations. Ranks are provided across the 56 possible combinations of regression models and embeddings (see Table S3). The best performance in each column is highlighted in bold.

achieved comparable Spearman's correlations (e.g., 0.840 for averaging vs. 0.835 for stacking, using eight top-performing combinations). While averaging performed slightly better when using fewer models (Figure 4b), ridge stacking exhibited a greater resistance to regression toward the mean, as indicated by a higher standard deviation ratio (SDR) (e.g., SDR = 0.81 for averaging vs. 0.88 for stacking, using eight top-performing combinations, Figure 4c). Due to the substantial increase in regression to the mean effect when averaging, ridge regression stacking was selected as the ensemble approach for the remainder of the study.

Performance consistently improved as the number of input predictions increased (Figure 4b and Table S4).

Notably, the ensemble model based on the top four different embeddings approached the highest Spearman's correlation ($\rho = 0.832$), even if the individual models that underpin it have lower performance than the four used in the corresponding top-performing ridge stacking (Figure 4b, c). Although stacking more than four top-performing models slightly improved the SDR (0.87 for six models vs. 0.86 for four different embeddings), the four-embedding approach was preferred for the remainder of the study due to substantially reduced computational demand at inference time.

Additional strategies to further improve ensemble learning were attempted by incorporating more information besides the

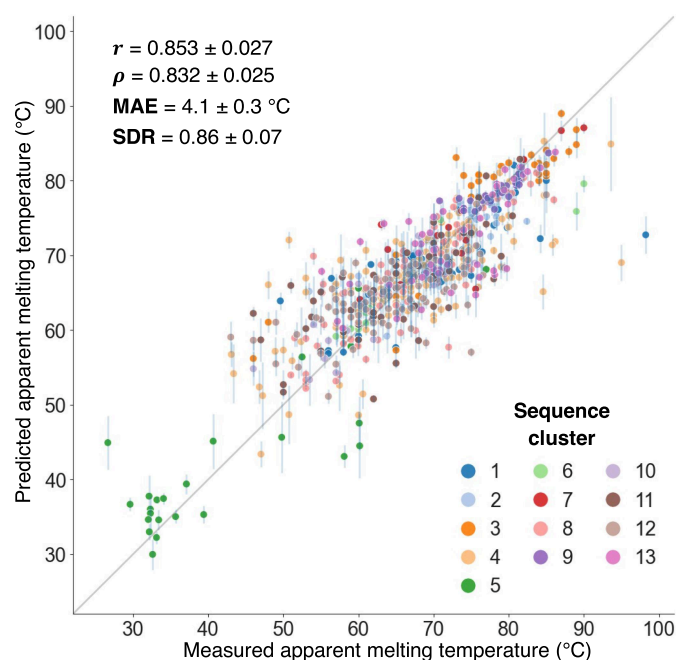


Figure 5. Test performance of the ensemble model across the nanobody dataset. Test-set prediction versus measured T_m from the final ensemble model, which consists in ridge regression stacking of models trained on diverse embeddings. For each nanobody, the prediction corresponds to the test prediction from the outer test fold of the nested cross-validation, averaged over three repeats. Error bars indicate the standard deviation of the predictions over the three pipeline repeats. The average standard deviation of the predictions over the repeats reaches 1.3°C. Scatter point colours represent the sequence k-medoids cluster of each nanobody. The reported Pearson's correlation (r), Spearman's correlation (ρ), MAE, and standard deviation ratio (SDR) are averaged across all outer test folds and repeats of the nested cross-validation, with corresponding standard deviations. The SDR assesses the model's tendency to regress towards the mean value of the dataset. Overall performance metrics for the plotted data are $r = 0.862$, $\rho = 0.845$, $\text{MAE} = 3.9^{\circ}\text{C}$, and $\text{SDR} = 0.84$.

predicted input temperatures, following approaches similar to the DeepSTABp framework.²⁶ Attempted augmentations included concatenating the input sequence embeddings or adding experimental conditions as input features (see Materials and Methods), but these did not significantly improve model performance (Table S5) and were thus not implemented in the final model.

This investigation culminated in an ensemble model that uses regression models trained on four diverse embeddings for nanobody thermostability prediction, which we call NanoMelt. Test-set predictions across the entire dataset of 640 datapoints, through various test folds of the outer cross-validation layer, are shown in Figure 5. The model demonstrates robust predictive performance, achieving a Pearson's correlation of 0.853, a Spearman's correlation of 0.832, a MAE of 4.1°C, and an SDR of 0.86.

To assess model behavior across different sequence clusters, we plotted individual cluster test-set performances (Fig. S9). High test performance was observed across both diverse and less diverse clusters. For instance, the highly diverse cluster 4 achieved a Pearson's correlation of 0.741, while the less diverse cluster 9, comprising point mutations from Sulea et al.,⁵⁵ reached 0.843. The consistency in test performance across clusters indicates that the model provides reliable predictions for both highly diverse nanobody sequences and point mutations of a given nanobody.

Finally, to set realistic expectations for model performance in real-world applications, we bootstrapped the Pearson's and

Spearman's correlations calculated across subsets of our dataset, sampling from 2% to 75% of sequence predictions in test-set folds (Fig. S10). This analysis addresses a common misconception: when applying a predictive model to new sequences, there is often an expectation that the correlation coefficients will match those reported during the original model assessment, which is typically based on much larger test datasets. Our results show that when sampling one-third of the sequences – the proportion used in the three-fold cross-validation – the correlation coefficients exhibit a narrow normal distribution centered around $r = 0.864$ and $\rho = 0.844$, as expected (Fig. S10). Similar patterns are observed when sampling 10% and 75% of the sequences. However, when evaluating performance on smaller subsets, such as a dozen sequences (2% sampling), the correlation coefficients display much broader distributions, with a 1.5 interquartile range spanning from $r = 0.597$ and $\rho = 0.499$ to nearly 1, and some outlier samples dropping almost to $r = 0$. These results underscore a critical point relevant to any data-driven model across various fields: reliable performance assessments require sufficiently large and diverse datasets (here approximately 50–60 sequences, as suggested in Fig. S10). In contrast, testing on small datasets inherently leads to substantial variability in observed performance, highlighting the unreliability of conclusions drawn from such limited assessments.

We also train an “operational” final model trained on the whole dataset to maximize data utilization (see Materials and Methods), which we make available as a webserver at <https://www-cohsoftware.ch.cam.ac.uk/> and for download from our repository (<https://gitlab.developers.cam.ac.uk/ch/sormanni/nanomelt>). NanoMelt takes around 250 sec to process 1,000 sequences on a GPU (NVIDIA RTX 8000).

Benchmarking

We benchmarked NanoMelt against three existing methods capable of predicting melting temperatures or of providing a stability score for protein sequences (Fig. S11). These methods include FoldX,⁵⁶ an empirical force-field-based algorithm that calculates the free energy of unfolding for a given structure; DeepSTABp,²⁶ a transformer-based melting temperature predictor trained on over 35,000 datapoints; and sequence likelihood outputs from ESM-2 and AntiBERTy, used as proxies for stability in the context of zero-shot predictions.

Given FoldX's reported sensitivity to minor atomistic displacements,^{6,56} we selected 46 nanobodies from our dataset that have corresponding crystal structures in the Protein Data Bank (PDB), to avoid biases due to inaccuracies in structural modeling. Our model achieved a Pearson's correlation of $r = 0.702$ on test-set predictions (i.e., when these 46 nanobodies were in the external test set). However, FoldX predictions showed no significant correlation with experimentally measured melting temperatures (Fig. S11a, $r = 0.033$), confirming that while empirical energy-based algorithms like FoldX are useful for predicting stability changes upon mutation, they struggle to assess stability across different but related protein sequences, such as nanobodies.⁵⁷

DeepSTABp was benchmarked on the entire dataset, but its predictions exhibited a weak Pearson's correlation of 0.267 and a very strong regression toward the mean of its training data (Fig. S11b).²⁶

We also tested the sequence likelihood from ESM-2 and the pseudo-log-likelihood from AntiBERTy as proxies for stability. Higher likelihood values are presumed to indicate more native-like sequences, which may correlate with greater stability, as suggested by Chungyoun et al.⁵⁸ However, these proxies did not show strong correlations with experimental data (Fig. S11c-d). Notably, ESM-2 achieved a Pearson's correlation of 0.337 across the entire dataset (p-value $\sim 10^{-18}$), outperforming stability-specific approaches like FoldX and DeepSTABp and confirming its potential for zero-shot prediction of protein fitness.³ However, further semi-supervised training or transfer learning for downstream prediction tasks appears essential to achieve accurate predictions.

Selection of highly stable nanobodies with NanoMelt

The most common approach to discovering new nanobodies targeting antigens of interest is through in vitro panning of VHH libraries, either obtained from immunized camels or constructed as naive libraries. Given the short sequence of nanobodies, next-generation sequencing (NGS) has become a standard practice for gaining a comprehensive view of the nanobody repertoire post-screening. Therefore, rapid, sequence-based computational predictors of biophysical traits, like NanoMelt, can significantly enhance these discovery pipelines by enabling the selection of high-fitness nanobodies for further experimental characterization.

To mimic this scenario and validate NanoMelt on independent data, we mined NGS-derived databases of camelid nanobody sequences⁵⁹ and selected six candidates for expression and characterization. Selection criteria included: 1) at least

30% sequence dissimilarity from the most similar sequence in our T_m nanobody dataset; 2) an AbNatiV VHH-nativeness score greater than 0.85; and 3) predicted thermostability: three nanobodies should be poorly stable (predicted $T_m < 61^\circ\text{C}$; Nb1, Nb2, Nb3) and three highly stable (predicted $T_m > 73^\circ\text{C}$; Nb4, Nb5, Nb6). AbNatiV,⁵⁹ a recently introduced deep learning framework, accurately predicts nanobody nativeness, which relates to the likelihood of a sequence belonging to the distribution of immune-derived nanobodies. We applied a stringent nativeness filter (>0.85) to ensure all selected sequences do not have highly unusual features that may correspond to sequencing or PCR-amplification errors, which are common in NGS datasets. The 30% minimal sequence dissimilarity to any sequence in our dataset further ensures a rigorous test of our method's ability to generalize to sequences very distant from those it was trained on. Sequences passing the first two filters show predicted T_m values between 58°C and 79°C (Fig. S12).

We transiently transfected mammalian cells for expression (see Materials and Methods) and found that all three nanobodies predicted to be highly stable expressed well, while none of the three predicted to be poorly stable showed any expression in the supernatant (Figure 6a-b and S13). We then purified the three expressing nanobodies, confirmed their identity by mass spectrometry (Table S6), and measured their T_m values. Remarkably, the measured values closely matched the predicted ones, with absolute errors of 1.7°C , 0.6°C , and 0.6°C for Nb4, Nb5, and Nb6, respectively (Figure 6a-b).

It is well established that low conformational stability can impair recombinant protein expression, often leading to very low yields or complete lack of soluble expression, as observed in our case. Thus, the lack of expression aligns well with the prediction of low T_m values for these nanobodies. We also considered thermodynamic solubility as another factor influencing expression yield. We ran CamSol intrinsic solubility

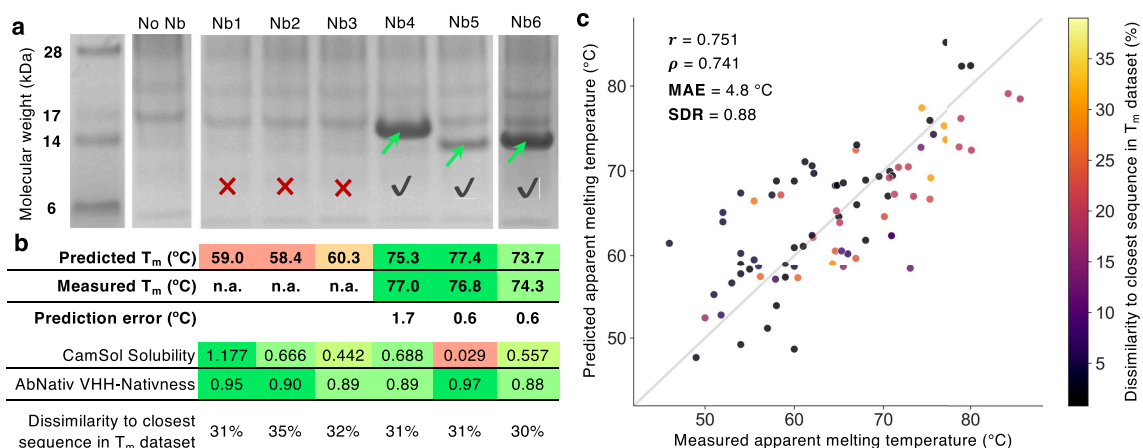


Figure 6. Real-world application of NanoMelt to sequences distant from those in T_m dataset. **(a)** SDS-page analysis of mammalian cell supernatants transiently transfected with six selected nanobodies (header). The “no Nb” lane represents the supernatant of cells not overexpressing a nanobody, serving as a negative control. The Nb6 sample was run on a separate gel (see fig. S12 for the uncropped images and additional results from repeated independent transfections). Green arrows indicate the nanobody bands at the expected molecular weight (confirmed by Mass spec, table S6), a red cross denotes lack of expression, and a grey tick indicates successful expression. **(b)** Table summarising the predicted T_m , measured T_m , prediction error, CamSol intrinsic solubility score, AbNatiV VHH-nativeness, and percentage dissimilarity from the closest sequence in the T_m dataset. Cells related to biophysical traits are colour-coded green, yellow, and red to represent favourable, intermediate, and unfavourable traits, respectively. **(c)** Prediction performances of NanoMelt on 83 external nanobody sequences (80 from various sources, plus the 3 that expressed from panel b). The Pearson's correlation (r), Spearman's correlation (ρ), MAE, and SDR are reported. Points are coloured based on sequence dissimilarity to the closest sequence in the training set. 31 sequences have a dissimilarity to the training set above 10%. Colouring based on the measurement technique is presented in figure S14. The grey line is the unity line.

predictions, which have consistently shown very high accuracy in ranking related proteins such as nanobodies in diverse settings.^{8,36,60,61} The results show that the non-expressing nanobodies had rather high predicted solubility, comparable to or higher than that of the two best-expressing nanobodies (Nb4 and Nb6; **Figure 6a-b**). Interestingly, Nb5, which has a much lower predicted solubility than all other nanobodies, also showed a visibly lower expression yield (**Figure 6a-b** and S12), supporting the idea that while conformational stability is a primary determinant of expression yield, solubility plays an important role, particularly when protein variants have similar stability.

After our initial study was completed, we collected 80 additional nanobody sequences and their melting temperatures from various sources (see Supplementary Dataset 2, Materials and Methods). NanoMelt achieves a Pearson's correlation of 0.751 and a MAE of 4.8°C on this external dataset, and the prediction error does not appear to depend on the dissimilarity to the closest sequence in the training set (**Figure 6c**).

Overall, this experimental validation on nanobody sequences that are highly divergent from those in the T_m dataset, along with the preserved accuracy on an external dataset of 80 sequences, reinforce our confidence in the practical utility of NanoMelt. NanoMelt has the potential to streamline nanobody discovery and optimization, offering a reliable tool for selecting highly stable, well-expressing candidates during the early stages of nanobody development.

Discussion

In this study, we conducted a comprehensive investigation into protein biophysical-trait learning using limited data and applied it to develop NanoMelt for nanobody thermostability prediction. We compiled a novel dataset of apparent melting temperatures, including 129 newly measured datapoints, resulting in a curated dataset of 640 labeled unique nanobody sequences (**Figure 3** and S5). By making this dataset publicly available, we aim to provide a valuable resource for the field (Supplementary Dataset 1).

As is common in machine learning with limited data, our dataset is somewhat noisy, with most datapoints aggregated from the literature and obtained using varying experimental protocols and data-analysis techniques. To address this variability, we applied stratification throughout the study to ensure consistent representation of experimental methods across all training and testing splits. However, we find that explicitly adding the experimental method as input to the model does not improve performance, possibly because the biggest source of variability comes from the usage of different buffers, heating rates, and protein concentrations, an information that is not available for most entries in the dataset. Evaluation on an external dataset confirmed that the experimental method is not a significant factor impacting performance (**Fig. S14**). For example, DSC-measured datapoints, which represent only 2.3% of the training dataset (**Figure 2b**), show low prediction errors (MAE = 3.7°C for 12 sequences of 83 in the external dataset).

We explored multiple-regression models with various sequence embeddings as input, utilizing repeated nested cross-validation to prevent data leakage between model tuning and performance assessment.¹⁰ Our analysis with artificial data highlighted the importance of input representation in protein fitness prediction. The performance gap between pre-trained embeddings and categorical encodings is more pronounced with smaller datasets and narrows as the dataset size increases (**Fig. S1**). While the preliminary study was performed with solubility-related predictions as a proxy for complex biophysical properties, similar dataset-size behaviors are observed when applying the same approach to a generated dataset of NanoMelt-predicted T_m (**Fig. S15**).

Interestingly, embeddings pre-trained specifically on antibodies and nanobodies did not outperform those pre-trained on generic proteins, despite being reported to be better at their primary task of sequence reconstruction.^{39–41} For instance, ESM-derived embeddings consistently outperformed antibody-specific embeddings, which in turn outperformed nanobody-specific embeddings, in predicting both nanobody calculated solubility scores (**Fig. S1**) and actual thermostability data (**Figure 4a**). This finding may reflect the broader applicability of the much bigger ESM embeddings and models, which are trained on larger and more diverse datasets, thus providing more comprehensive information relevant to general physicochemical properties of polypeptides.⁶² We further find that later layers in pre-trained models yielded the most effective embeddings, with the performance of the last layer being at par with that of the second-last layer (**Fig. S16**).

Pre-trained embeddings also reduce overfitting, which is consistently greater at small training set sizes (**Fig. S1**) and is thus an important concern when learning from limited data, as it can hinder generalization to unseen data.⁶³ In this study, we systematically reported both training and testing performance to assess the tendency of various regression models to overfit. This practice, which is often overlooked in recent literature,¹⁴ ensures a transparent assessment of model performance and its likelihood to generalize to more distant sequences. To mitigate overfitting, feature selection was applied to each regression model to reduce the number of free parameters. In our ensemble model, we used a diverse range of input embeddings and regressors, rather than stacking only the best-performing combinations of embedding and models, to mitigate the dependency on individual models' biases. Moreover, only out-of-fold predictions were used as inputs when training the stacking ridge model, so it could learn to leverage the individual predictions in a more realistic, inference-like scenario.

A key challenge with small datasets is the tendency for models to regress toward the mean, which can result in promising correlation coefficients despite a lack of coverage across the full range of measured values.⁶⁴ We addressed this by calculating the SDR, which quantifies this tendency. For example, while simple averaging of predictions performed at par or better than ridge regression stacking in terms of Spearman's correlation, the SDR revealed a stronger tendency to regress toward the mean with averaging, leading us to prefer ridge regression for the final model (**Figure 4b**).

While NanoMelt's MAE is 4.1°C (Figure 5), some nanobodies exhibited much larger prediction errors, around 20°C. These outliers often had unusually high or low melting temperatures compared to similar sequences in the training set. For instance, the W37A mutant of a nanobody with a wildtype (WT) melting temperature of 60.2°C had an apparent melting temperature of 26.6°C, but its prediction was 47.5°C. Although the model correctly indicated strong destabilization, it underestimated the extent of the effect, likely due to the mutation involving a hyper-conserved site.

In conclusion, we have introduced a computational pipeline that combines multiple features to enable accurate learning from small datasets (Figure 1). This pipeline has the potential to be applied to the prediction of other properties and across different protein classes, including complete Fv antibodies. However, the VH/VL pairing in Fv antibodies may introduce additional complexity, likely necessitating a larger dataset. Nanobodies are crucial reagents in research, diagnostics, and therapeutics, and NanoMelt can significantly streamline their development by providing a reliable and rapid tool for selecting highly stable candidates during discovery and optimization campaigns.

Materials and methods

Input sequence representation

Eight distinct embeddings were studied to represent the sequences as inputs for the regression models (see Table S1). The categorical one-hot encoding is constructed by aligning each sequence on the 149-long AHo numbering scheme⁶⁵ via the ANARCI software.⁶⁶ Each position is represented by a binary vector of size 21 filled with 0 except for a single 1 at the alphabet index of the residue under scrutiny (20 standard amino acids and a gap position). These per-position vectors are concatenated together to lead to a one-hot vector of 3,219 elements. Similarly, the physicochemical-based VHSE encoding characterizes each residue with vectors of size 8 derived from the principal component analysis previously conducted on 50 hydrophobic, steric, and electrostatic descriptors.³⁷ Gap residues in AHo-aligned sequences are represented here with 8 zeros. The per-position vectors are concatenated together to lead to a VHSE vector comprising 1,192 elements.

Two embeddings are respectively generated by the large language protein models ESM-1b³⁸ (650 M parameters) and ESM-2⁴ (150 M parameters). Both models train a transformer with an unsupervised masked language objective on sequences from a wide range of protein classes from the Universal Protein Knowledgebase (UniProt).⁶⁷ The ESM-2 model is an updated version of the ESM-1b model with a bigger training dataset (65 M and 30 M sequences, respectively), and distinct architecture features (e.g., type of positional embedding). Both the ESM-1b and ESM-2 embeddings are extracted from the last layer of their transformer (the 33rd and 30th layers, respectively). The per-residue representations are averaged across all positions, excluding the beginning- and ending-of-sequence tokens, yielding 1,280- and 640-sized embeddings for ESM-1b and ESM-2, respectively.

Three additional embeddings are generated by antibody-specific protein language transformers. AntiBERTy⁴⁰ (26 M

parameters) is a transformer trained on 42 M heavy antibody sequences from the Observed Antibody Space (OAS) database⁶⁸ using a multiple instance learning framework to predict the likelihood of residues in the sequences. A per-residue representation is extracted from its last layer and averaged across all positions, excluding the beginning- and end-of-sequence tokens, leading to a 512-sized embedding. The AbLang model³⁹ (8.5 M parameters) is trained on 14 M heavy antibody sequences from the OAS database with a masked language reconstruction objective. A per-residue representation is extracted from its last layer and averaged across all positions resulting in a 768-long embedding. The nanoBERT model⁴¹ (14 M parameters) follows the training protocol of AntiBERTy, but it is trained specifically on 10 M nanobody sequences from the Integrated Nanobody Database for Immunoinformatics (INDI) database.⁶⁹ A per-residue representation is extracted from its last layer and averaged across all positions, excluding the beginning- and end-of-sequence tokens, to lead to a 320-long embedding.

Lastly, the nanobody structure predictor NanobodyBuilder2⁴² is also used to generate structure-aware embeddings. NanobodyBuilder2 is a set of four transformer models trained on ~2k nanobody structures from the structural antibody database (SAbDab).⁷⁰ A per-residue representation is extracted from the last layer of the first model (7.6 M parameters) and averaged across all positions yielding a 128-long embedding.

Minimal dataset size search

Ten thousand nanobody sequences were randomly sampled from the training dataset of native camelid nanobodies of the AbNatiV nanobody nativeness model.⁵⁹ Their respective structure was modeled by NanoBodyBuilder2.⁴² For each nanobody, the intrinsic solubility and structurally corrected scores were computed using the CamSol algorithm.³⁶ The intrinsic score is solely computed on the amino acid sequence, while the structurally corrected score corrects for residue solvent exposure and proximity in the modeled structure.

Both generated artificial solubility datasets were randomly partitioned with an 80:20 ratio. From the 8,000 sequences, we sampled 14 subsets of varying size ranging from 20 to 8,000 datapoints. We trained a ridge regression on each subset with hyperparameters tuned via a 5-fold cross validation by grid-search, as implemented in Scikit-learn.⁷¹ For each trained model, we computed the Spearman's coefficient to assess the prediction performance on the held-out test set of 2,000 sequences. The process was repeated over five random splitting seeds, and the standard deviation of the performance across the five splits is plotted as error bars (Fig. S1).

Expression, purification, and preparation of nanobodies for T_m dataset

Ninety-three nanobody sequences were selected by cluster sampling from the PDB⁷² to ensure diversity. The PDB was used as a source because most nanobodies therein come from different animals immunized at different time, hence providing a diverse sample, and – more importantly – because all nanobodies therein are known to have recombinantly expressed successfully and have a sequence free of any

sequencing errors (which are common in NGS datasets). The selected sequences were expressed using the turbo-CHO platform (GenScript) in mammalian cells with a C-terminal His6-tag and purified via immobilized metal affinity chromatography in phosphate-buffered saline (PBS) buffer. Purity was confirmed by SDS – polyacrylamide gel electrophoresis (SDS-PAGE) analysis. All protein concentrations were adjusted between 0.10 and 0.18 mg.ml⁻¹ to reduce aggregation upon heating while retaining a good signal-to-noise. Nanobody concentrations were determined using blanked absorbance 280 nm values in duplicates with the Nanodrop ND-1000 instrument (PepLab Biotechnologie, Erlangen, Germany). The extinction coefficients were calculated from the amino acid sequence using the ExPASy ProtParam tool (web.expasy.org/protparam/).

Nano-differential scanning fluorimetry assays

NanoDSF measurements were conducted on a Prometheus NT.48 instrument (NanoTemper Technologies, Munich, Germany) with high-sensitivity capillaries. Fluorescence was monitored at 330 nm and 350 nm with an excitation wavelength of 280 nm. Temperature was increased at a rate of 2 °C/min starting from 22°C and ending at 95°C. All measurements were performed in duplicates in different runs using different capillary positions to mitigate possible systematic errors.

Melting curve fitting

For the 191 nanoDSF curves analyzed in-house, the individual fluorescence signals at 350 nm and 330 nm, and the 350/330 ratio trace were processed to determine a representative melting temperature per nanobody. Each fluorescence trace was first smoothed via a Savitzky-Golay filter (window length = 21, polynomial order = 2) and fitted via least-squares minimization on the following two-state thermal denaturation model⁷³:

$$y = \frac{\alpha_N + \beta_N T + (\alpha_D + \beta_D T) \exp\left(\frac{\Delta H_{D-N}}{R} \left(\frac{1}{T_M} - \frac{1}{T}\right)\right)}{1 + \exp\left(\frac{\Delta H_{D-N}}{R} \left(\frac{1}{T_M} - \frac{1}{T}\right)\right)} \quad (1)$$

where y is the measured signal, α_N , β_N and α_D , β_D the intercept and slope of the linear baselines of the native and denatured states respectively, R the molar gas constant, ΔH_{D-N} the enthalpy of equilibrium between the denatured and native states, and T_M the apparent melting temperature.

A T_M of 65°C, and a ratio $\frac{\Delta H_{D-N}}{R}$ of 3000 °C were given as initial parameter guesses. Linear fitting is done at the beginning and ending of the signal to obtain the initial guesses of the slope and intercept of the native and denatured states. For the individual 350 nm and 330 nm fluorescence signals, the traces were initially corrected by subtracting it with its own fitted linear regression to provide more consistent data to the fitting algorithm, thus enhancing the transition regime (a coarse but good-enough correction for the temperature-dependence of the intrinsic fluorescence). For the individual 350 nm and 330 nm signals, the final fitting is only carried on

a window of 10°C around the T_M extracted from the fitting of the 350/330 ratio trace beforehand.

For each analyzed nanobody, the representative apparent melting temperature was calculated by averaging the T_M values extracted from the fits of the individual 350 nm and 330 nm fluorescence signals, whenever conclusive. Each fit was visually inspected for accuracy; if deemed unreliable, the T_M value extracted from the 350/330 ratio fit was taken instead. Among the 191 nanobodies analyzed, the T_M values for 166 nanobodies were extracted from the individual signals, and for 25 nanobodies from the ratio, as the unfolding transition was too subtle for reliable fitting in the individual signals.

Construction of the nanobody thermostability dataset

A dataset of 640 nanobody melting temperatures was built in this study. The datapoints are derived from three main sources (see Table S2):

- 193 nanobody sequences from our own curated dataset, including 95 nanobodies selected from the PDB database, 34 nanobodies already available in our laboratory from previous projects, and raw nanoDSF data for 64 additional nanobodies from the research conducted by Kunz et al.⁴⁴. 68 nanobodies were published in the latter work, yet 3 could not be aligned to the AHo scheme, and the signal of NbPep53 was too poor for a T_m to be extracted (T_m likely above 90°C with no lower plateau observed). All measurements were performed under the same experimental conditions with nanoDSF and analyzed by fitting two-state thermal denaturation model to ensure high consistency. As a control, six nanobodies from the publication of Kunz et al.⁴⁴ were produced and characterized along our selected nanobodies. The T_m we obtained align closely with the data reported in the original publication of Kunz et al. (Fig. S2).
- 405 nanobodies from the NbThermo database.³⁵ These datapoints were collated from various previously published work, and the melting temperatures were measured with a range of experimental methods, including nanoDSF, DSF, CD, and DSC. Most of the melting temperatures reported were determined by steepest derivative analysis. The original NbThermo database of 548 nanobodies (json file downloaded in April 2023) was filtered following the alignment cleaning procedure in AbNatiV.⁵⁹ In summary, some sequences were removed because:
 - a. 29 entries did not have an associated sequence,
 - b. 6 entries did not have a melting temperature associated,
 - c. 9 sequences could not be aligned,
 - d. 63 sequences matching the ones from Kunz et al.⁴⁴ were removed as we re-analyzed them,
 - e. 18 entries had a duplicate sequence with another entry,
 - f. 18 others were duplicates of 14 sequences from our own dataset of 127 characterized nanobodies, particularly those expressed in previous projects and already published.

Furthermore, 50 additional sequences have multiple melting temperatures reported for different experimental methods. The 82 unique sequences with multiple values reported (i.e., sequences duplicates or multiple experimental methods) showed a mean pairwise absolute difference of 2.4°C, which can serve as an indication of the lowest MAE that may be expected from predictors trained on this dataset (Fig. S7). For these sequences, the retained melting temperature corresponds to the value obtained from the experimental method when available in the following order of preference: nanoDSF, DSF, DSC, CD, and others, giving preference to our own characterization when available.

- 42 nanobody sequences from mutational studies. Specifically, 23 datapoints come from the study by Rosace et al.,⁸ comprising three WT nanobodies with up to 4 mutations. The additional 19 datapoints originate from an unpublished destabilization study on a WT nanobody biosensor and were kindly provided to us by Dr Gabriel Ortega Quintanilla, who gave us permission to publish sequences and melting temperatures.

We make the whole dataset available in Supplementary Dataset 1.

Sequence clustering

Nanobody sequences of the curated thermostability dataset were clustered by the k-medoids algorithm from the Scikit-learn package.⁷¹ These sequences were one-hot embedded following AHO-numbering alignment prior to clustering. To determine the optimal number of clusters, the Kneedle algorithm⁷⁴ was used to identify an elbow cutoff point. Increasing the number of clusters beyond this elbow point does not result in a significant reduction of the K-medoids' inertia.

Training of the regression models

Seven regression models were tested in this study, ranging from linear models (i.e., Ridge,⁴⁶ Huber,⁴⁷ and Elastic Net⁴⁸) to non-linear ones (i.e., the decision tree-based algorithms Random Forest⁴⁹ and LightGBM,⁵⁰ Support Vector Machine,⁵¹ and Gaussian Process⁵²). All the ranges of hyperparameter ranges given during the tuning process are directly available in the code repository at <https://gitlab.developers.cam.ac.uk/ch/sormanni/nanomelt>.

Each model was trained and evaluated using a repeated nested stratified cross-validation procedure,¹¹ as illustrated in Figure 1. Every dataset splitting was stratified according to the type of experimental method used and the associated sequence similarity cluster. Stratification maintains the distribution of classes across the different folds during the splitting process.

The regression models were trained and ranked via a repeated nested three-fold cross-validation. This procedure

incorporates two layers of cross-validation: an outer loop and an inner loop. In the outer loop, the dataset was split into three folds. Two folds were used by the inner loop to train and tune the hyperparameters on a three-fold grid search, while the remaining fold served as the test set to assess the performance of the tuned models. To enhance the reliability of the results, the entire nested procedure was repeated with three different random seeds, and the performances of each regression on the test sets were averaged. During hyperparameter tuning, each inner loop was also repeated three times, and the performances were averaged across these repeats to select the best model.

To ensure that regularization is applied uniformly across the features (as for the Ridge, Elastic Net, SVR and GPR models), we standardized the embeddings within each cross-validation by removing the mean of the training subset and scaling it to unit variance. In addition, to reduce overfitting, we performed feature selection based on importance rankings obtained from a Ridge regression model trained on the corresponding training subset at each fold.

Model selection, feature selection and ensemble learning

The predicted temperatures derived from the individually trained regression models served as inputs to train the ensemble model. These predicted temperatures were collected from each outer test fold during the nested cross-validation of every model. Three sets of predicted temperatures were collected from the three repeats of the pipeline on three different seeds. An ensemble model was trained following the same stratified nested cross-validation pipeline on each set of predicted temperatures with its associated seed. The performances of the ensemble model were averaged across the three repeats.

For ensemble learning, three model selection strategies were explored to combine the various regression models. The first one selects up to the eight top ranking regression models, including the type of embedding, from the ranking of the repeated stratified nested cross-validation evaluation (Table S3). The second one selects the best performing regression model for each embedding type, with up to the eight embeddings (Table 1). And the third one selects the best performing embedding for each regression model, with up to the seven models (Table S3).

Three feature selection strategies were explored as inputs for the ensemble ridge model. The first one added the sequence information by concatenating the ESM-1b embedding or one-hot encoding with the eight predicted input temperatures (row A in Table S5). Either the size of the original embedding was maintained, or it was reduced with a linear layer to match the number of input temperatures. The second strategy added the experimental method information by concatenating it to the predicted input temperatures (row B in Table S5). Either the method information was added as a single integer (e.g., 0 for nanoDSF and 1 for DSF), or expanded via a linear layer to match the number of input temperatures. The third strategy integrated both sequence and experimental information into the ensemble model (row C in Table S5).

A ridge regression model was used as the stacking model for ensemble learning. It was trained following the same repeated stratified nested cross-validation procedure. The stacking model takes as inputs the predicted melting temperatures from the outer test folds of the previous nested cross-validation for each trained regression model.

Deployment of the final thermostability predictive model

For final model deployment, each individual regression model was tuned first using the stratified cross-validation but then re-trained on the whole dataset. Subsequently, the stacking ridge regression and error predicting GPR were trained and re-trained following the same procedure, always using as input the predicted melting temperatures from the test folds of the previous individual cross-validations.

Expression, purification, and characterization of the six nanobodies of Figure 6a.

Amino acid sequences encoding selected nanobodies were codon-optimized for human-cell expression using the Genscript online tool. Resulting DNA sequences were ordered from Genscript as gene fragments flanked by Golden Gate cloning sites for insertion in a pcDNA3.4 mammalian expression vector (Addgene), modified to harbour a CD33 signal secretion sequence at the N-terminus and a 6×His tag at the C-terminus of the cloning sites. Gene fragments were cloned using Golden Gate cloning with BsmBI-v2 golden gate assembly kit (New England Biolabs; E1602S). Resulting plasmids were transformed into Dh5α competent *E. coli* cells (ThermoFisher Scientific) and grown overnight at 37°C on LB media plates containing ampicillin, before midi prep cultures were set up the next day. Midi preps were processed using QIAGEN Midi Prep kit (QIAGEN). Purified plasmids were sent for sequencing and, upon confirmation of the correct sequence, were used for transfection.

Plasmids were transfected into Expi293F cell line following protein transfection protocol from the Expi293 transfection kit manufacturer (ThermoFisher Scientific; A14635). 3 ml cultures were set up for each nanobody. Cells were incubated for 5 days at 37°C with 5% CO₂ on an orbital shaker with 120 rpm. On day 5, cells were harvested by centrifugation and the supernatant was collected for SDS-PAGE analysis and subsequent protein purification.

His Mag Sepharose Excel magnetic beads (Cytiva) were washed with PBS before being added to protein supernatants. For each 3 mL culture, 100 µL of beads were added and samples were incubated on a roller at 4°C for 2–3 hours. Beads were then washed and resuspended with PBS to be processed on AmMag™ SA Plus Semi-automated purification System 980 (Genscript). Beads were washed with PBS at 4 mM imidazole and protein eluted in PBS at 200 mM imidazole. Eluted proteins were then purified further by size-exclusion chromatography on an AKTA system, to remove the imidazole and to further purify the monomeric nanobodies. An AKTA pure 25 system was used with a Superdex 75 increase 10/300 GL column with PBS as a running buffer. Resulting purified proteins in PBS were aliquoted and flash frozen in liquid nitrogen. Samples were then

stored in –80°C. The mass of each protein was confirmed by comparing the theoretical molecular weight (from the ExPASy ProtParam tool) to that obtained from liquid-chromatography mass spectroscopy using VION (Waters).

The apparent melting temperature of the three nanobodies that expressed was measured with a Tycho NT.6 (Nanotemper). Intrinsic fluorescence of tryptophan and tyrosine residues was recorded at 330 and 350 nm with a fixed 30 °C/min temperature ramp from 35 to 95°C. Data were then analysed as described in the “Melting curve fitting” section above to obtain T_m values.

External thermostability dataset with 83 nanobodies

Since the initial submission of this work, 83 additional nanobody sequences and their melting temperatures have been collected from various sources (Supplementary Dataset 2):

- 10 sequences were expressed and characterized internally following the same protocol described previously, these include the 3 of Fig. 6a-b.
- 59 sequences were assembled from multiple sources of the literature by Alvarez *et al.*⁷⁵ The experimental methods used were retrieved from the original sources.
- 11 sequences were characterized by DSC by Gómez-Mulas *et al.*^{76,77}
- 3 sequences characterized by CD were collected from an internal thesis.

The distance to NanoMelt’s training set is reported in Figure 6.b and the distribution of experimental methods in Figure S14. This dataset was used solely for external evaluation.

Acknowledgments

We are grateful to Dr Chris Johnson for his support in facilitating the use of the Prometheus equipment at the MRC Laboratory of Molecular Biology. We acknowledge Dr Gabriel Ortega Quintanilla for sharing his thermostability data from an unpublished study. We thank Montader Ali, Misha Atkinson, Magdalena Nowinska, Dr Mauricio Aguilar Rangel, and Dr Oded Rimon for donating samples of their purified nanobodies to expand our dataset. P.S. is a Royal Society University Research Fellow (grant no. URF\R1\201461). We acknowledge funding from UK Research and Innovation (UKRI) Engineering and Physical Sciences Research Council (grant no. EP/X024733/1, an ERC starting grant to P. S. underwritten by UKRI).

Disclosure statement

No potential conflict of interest was reported by the author(s).

Funding

The work was supported by the Engineering and Physical Sciences Research Council [EP/X024733/1]; Royal Society [URF\R1\201461].

ORCID

Aubin Ramon  <http://orcid.org/0009-0002-5502-5961>
Pietro Sormanni  <http://orcid.org/0000-0002-6228-2221>

Author contributions

PS conceived the project with PK. PS supervised the project and secured funding. AR developed the machine learning pipeline with the help of MN. AR measured and analyzed the nanoDSF curves. AR parsed and collected the nanobody thermostability data. OP and RG produced nanobodies and carried out wet-lab experiments. AR and PS wrote the first version of the paper. All authors analyzed data and edited the paper.

Data availability statement

The NanoMelt repository is accessible at <https://gitlab.developers.cam.ac.uk/ch/sormanni/nanomelt>. It includes all the data used in the study, the NanoMelt trained model, the python pipeline, and the notebooks used for analysis. A user-friendly webserver to run NanoMelt is provided at <https://www-cohsoftware.ch.cam.ac.uk/>. To access the webserver, users need to register a free account and log in

Abbreviations

GPR	Gaussian Process Regression
T _m	Apparent Melting Temperature
CD	Circular Dichroism
DSC	Differential Scanning Calorimetry
DSF	Differential Scanning Fluorimetry
GPR	Gaussian Process Regression
MAE	Mean Absolute Error
mAb	Monoclonal Antibody
NGS	Next-Generation Sequencing
PBS	Phosphate-Buffered Saline
PDB	Protein Data Bank
SDS-PAGE	SDS – Polyacrylamide Gel Electrophoresis
SDR	Standard Deviation Ratio
SAbDab	Structural Antibody Database
SVR	Support Vector Regression
WT	Wildtype

References

- Molina RS, Rix G, Mengiste AA, Álvarez B, Seo D, Chen H, Hurtado JE, Zhang Q, García-García JD, Heins ZJ, et al. In vivo hypermutation and continuous evolution. *Nat Rev Methods Primer*. 2022;2(1):1–22. doi:10.1038/s43586-022-00119-5.
- Kim JY, Yoo H-W, Lee P-G, Lee S-G, Seo J-H, Kim B-G. In vivo protein evolution, next generation protein engineering strategy: from random approach to target-specific approach. *Biotechnol Bioprocess Eng*. 2019;24(1):85–94. doi:10.1007/s12257-018-0394-2.
- Meier J, Rao R R, Verkuil R, Liu J, Sercu T, Rives A. Language models enable zero-shot prediction of the effects of mutations on protein function in. In: Ranzato M, Beygelzimer A, Dauphin Y, Liang PS, Wortman Vaughan J, editors. *Advances in neural information processing systems*. Red Hook, NY, USA: Curran Associates, Inc.; 2021. p. 29287–29303.
- Lin Z, Akin H, Rao R, Hie B, Zhu Z, Lu W, Smetanin N, Verkuil R, Kabeli O, Shmueli Y, et al. Evolutionary-scale prediction of atomic-level protein structure with a language model. *Science*. 2023;379(6637):1123–1130. doi:10.1126/science.ade2574.
- Jumper J, Evans R, Pritzel A, Green T, Figurnov M, Ronneberger O, Tunyasuvunakool K, Bates R, Židek A, Potapenko A, et al. Highly accurate protein structure prediction with AlphaFold. *Nature*. 2021;596(7873):583–589. doi:10.1038/s41586-021-03819-2.
- Buř O, Rudat J, Ochsenreither K. FoldX as protein engineering tool: better than random based approaches? *Comput Struct Biotechnol J*. 2018;16:25–33. doi:10.1016/j.csbj.2018.01.002.
- Laurie ATR, Jackson RM. Q-SiteFinder: an energy-based method for the prediction of protein–ligand binding sites. *Bioinformatics*. 2005;21(9):1908–1916. doi:10.1093/bioinformatics/bti315.
- Rosace A, Bennett A, Oeller M, Mortensen MM, Sakhnini L, Lorenzen N, Poulsen C, Sormanni P. Automated optimisation of solubility and conformational stability of antibodies and proteins. *Nat Commun*. 2023;14(1):1937. doi:10.1038/s41467-023-37668-6.
- Barbero-Aparicio JA, Olivares-Gil A, Rodríguez JJ, García-Orsorio C, Díez-Pastor JF. Addressing data scarcity in protein fitness landscape analysis: a study on semi-supervised and deep transfer learning techniques. *Inf Fusion*. 2024;102:102035. doi:10.1016/j.inffus.2023.102035.
- Varma S, Simon R. Bias in error estimation when using cross-validation for model selection. *BMC Bioinf*. 2006;7(1):91. doi:10.1186/1471-2105-7-91.
- Krstajic D, Buturovic LJ, Leahy DE, Thomas S. Cross-validation pitfalls when selecting and assessing regression and classification models. *J Cheminf*. 2014;6(1). doi:10.1186/1758-2946-6-10.
- Safonova A, Ghazaryan G, Stiller S, Main-Knorn M, Nendel C, Ryo M. Ten Deep learning techniques to address small data problems with remote sensing. *Int J Appl Earth Obs Geoinf*. 2023;125:103569. doi:10.1016/j.jag.2023.103569.
- Xu P, Ji X, Li M, Lu W. Small data machine learning in materials science. *Npj Comput Mater*. 2023;9(1):1–15. doi:10.1038/s41524-023-01000-z.
- Hsu C, Nisonoff H, Fannjiang C, Listgarten J. Learning protein fitness models from evolutionary and assay-labeled data. *Nat Biotechnol*. 2022;40(7):1114–1122. doi:10.1038/s41587-021-01146-5.
- Ma J, Fong SH, Luo Y, Bakkenist CJ, Shen JP, Mourragui S, Wessels LFA, Hafner M, Sharan R, Peng J, et al. Few-shot learning creates predictive models of drug response that translate from high-throughput screens to individual patients. *Nat Cancer*. 2021;2(2):233–244. doi:10.1038/s43018-020-00169-2.
- Mardikoraem M, Woldring D. Protein fitness prediction is impacted by the interplay of language models, ensemble learning, and sampling methods. *Pharmaceutics*. 2023;15(5):1337. doi:10.3390/pharmaceutics15051337.
- Shahhosseini M, Hu G, Pham H. Optimizing ensemble weights and hyperparameters of machine learning models for regression problems. [Preprint] (2019). Accessed 2024 Sep 4.] doi:<http://arxiv.org/abs/1908.05287>.
- Kumar S, Tsai C-J, Nussinov R. Factors enhancing protein thermostability. *Protein Eng Des Sel*. 2000;13(3):179–191. doi:10.1093/protein/13.3.179.
- Liu Y, Hoppenbrouwers T, Wang Y, Xie Y, Wei X, Zhang H, Du G, Imam KMSU, Wichers H, Li Z, et al. Glycosylation contributes to thermostability and proteolytic resistance of rFIP-nha (Nectria haematococca). *Mol Basel Switz*. 2023;28(17):6386. doi:10.3390/molecules28176386.
- Mezzasalma TM, Kranz JK, Chan W, Struble GT, Schalk-Hihi C, Deckman IC, Springer BA, Todd MJ. Enhancing recombinant protein quality and yield by protein stability profiling. *J Biomol Screen*. 2007;12(3):418–428. doi:10.1177/1087057106297984.
- Vermeer AWP, Norde W. The thermal stability of immunoglobulin: unfolding and aggregation of a multi-domain protein. *Biophys J*. 2000;78(1):394–404. doi:10.1016/S0006-3495(00)76602-1.
- Kwon WS, Da Silva NA, Kellis JT. Relationship between thermal stability, degradation rate and expression yield of barnase variants in the periplasm of Escherichia coli. *Protein Eng Des Sel*. 1996;9(12):1197–1202. doi:10.1093/protein/9.12.1197.
- Timr S, Madern D, Sterpone F. Chapter six - protein thermal stability. In: Strodel B, and Barz B. editors. *Progress in molecular biology and translational science, computational approaches for understanding dynamical systems: protein folding and assembly*. Amsterdam, The Netherlands: Elsevier B.V.; 2020. p. 239–272.
- Nikam R, Kulandaisamy A, Harini K, Sharma D, Gromiha MM. ProThermDB: thermodynamic database for proteins and mutants revisited after 15 years. *Nucleic Acids Res*. 2021;49(D1):D420–D424. doi:10.1093/nar/gkaa1035.

25. Rollins ZA, Widatalla T, Cheng AC, Metwally E. AbMelt: learning antibody thermostability from molecular dynamics. *Biophys J*. 2024;(17):2921–2933. doi:10.1016/j.bpj.2024.06.003.
26. Jung F, Frey K, Zimmer D, Mühlhaus T. DeepSTABp: a deep learning approach for the prediction of thermal protein stability. *Int J Mol Sci*. 2023;24(8):7444. doi:10.3390/ijms24087444.
27. Yang Y, Zhao J, Zeng L, Vihinen M. ProTstab2 for prediction of protein thermal stabilities. *Int J Mol Sci*. 2022;23(18):10798. doi:10.3390/ijms231810798.
28. Zhao J, Yan W, Yang Y. DeepTP: a deep learning Model for thermophilic protein prediction. *Int J Mol Sci*. 2023;24(3):2217. doi:10.3390/ijms24032217.
29. Harmalkar A, Rao R, Richard Xie Y, Honer J, Deisting W, Anlahr J, Hoenig A, Czwikla J, Sienz-Widmann E, Rau D, et al. Toward generalizable prediction of antibody thermostability using machine learning on sequence and structure features. *mAbs*. 2023;15(1). doi:10.1080/19420862.2022.2163584.
30. Muyldermans S. Nanobodies: natural single-domain antibodies. *Annu Rev Biochem*. 2013;82(1):775–797. doi:10.1146/annurev-biochem-063011-092449.
31. Hamers-Casterman C, Atarhouch T, Muyldermans S, Robinson G, Hammers C, Songa EB, Bendahman N, Hammers R. Naturally occurring antibodies devoid of light chains. *Nature*. 1993;363(6428):446–448. doi:10.1038/363446a0.
32. Pardon E, Laeremans T, Triest S, Rasmussen SGF, Wohlkönig A, Ruf A, Muyldermans S, Hol WGJ, Kobilka BK, Steyaert J, et al. A general protocol for the generation of nanobodies for structural biology. *Nat Protoc*. 2014;9(3):674–693. doi:10.1038/nprot.2014.039.
33. Virant D, Traenkle B, Maier J, Kaiser PD, Bodenhöfer M, Schmees C, Vojnovic I, Pisak-Lukáts B, Endesfelder U, Rothbauer U, et al. A peptide tag-specific nanobody enables high-quality labeling for dSTORM imaging. *Nat Commun*. 2018;9(1):1–14. doi:10.1038/s41467-018-03191-2.
34. Carter PJ, Rajpal A. Designing antibodies as therapeutics. *Cell*. 2022;185(15):2789–2805. doi:10.1016/j.cell.2022.05.029.
35. Valdés-Tresanco MS, Valdés-Tresanco ME, Molina-Abad E, Moreno E. NbThermo: a new thermostability database for nanobodies. *Database*. 2023;2023. doi:10.1093/database/baad021.
36. Sormanni P, Aprile FA, Vendruscolo M. The CamSol method of rational design of protein mutants with enhanced solubility. *J Mol Biol*. 2015;427(2):478–490. doi:10.1016/j.jmb.2014.09.026.
37. Mei H, Liao ZH, Zhou Y, Li SZ. A new set of amino acid descriptors and its application in peptide QSARs. *Biopolymers*. 2005;80(6):775–786. doi:10.1002/bip.20296.
38. Rives A, Meier J, Sercu T, Goyal S, Lin Z, Liu J, Guo D, Ott M, Zitnick CL, Ma J, et al. Biological structure and function emerge from scaling unsupervised learning to 250 million protein sequences. *Proc Natl Acad Sci USA*. 2021;118(15). doi:10.1073/pnas.2016239118.
39. Olsen TH, Moal IH, Deane CM. AbLang: an antibody language model for completing antibody sequences. *Bioinform Adv*. 2022;2(1).
40. Ruffolo JA, Gray JJ, Sulam J. Deciphering antibody affinity maturation with language models and weakly supervised learning. *Machine Learning for Structural Biology Workshop, NeurIPS 2021*.
41. Thorling J, Hadsund JT, Satława T, Janusz B, Shan L, Röttger R, Krawczyk K. nanoBERT: a deep learning model for gene agnostic navigation of the nanobody mutational space. 2024. doi:10.1101/2024.01.31.578143.
42. Abanades B, Wong WK, Boyles F, Georges G, Bujotzek A, Deane CM. ImmuneBuilder: deep-learning models for predicting the structures of immune proteins. *Commun Biol*. 2023;6(1). doi:10.1038/s42003-023-04927-7.
43. Biswas S, Khimulya G, Alley EC, Esvelt KM, Church GM. Low-N protein engineering with data-efficient deep learning. *Nat Methods*. 2021;18(4):389–396. doi:10.1038/s41592-021-01100-y.
44. Kunz P, Zinner K, Mücke N, Bartoschik T, Muyldermans S, Hoheisel JD. The structural basis of nanobody unfolding reversibility and thermoresistance. *Sci Rep*. 2018;8(1):1–10. doi:10.1038/s41598-018-26338-z.
45. Žoldák G, Jancura D, Sedlák E. The fluorescence intensities ratio is not a reliable parameter for evaluation of protein unfolding transitions. *Protein Sci*. 2017;26(6):1236–1239. doi:10.1002/pro.3170.
46. Hoerl AE, Kennard RW. *Technometrics*. 2000;42(1):80–86.
47. Huber PJ. Robust estimation of a location parameter. In: Kotz S, and Johnson NL, editors. *Breakthroughs in statistics: methodology and distribution*. Beachwood, OH: Springer; 1992. p. 492–518.
48. Zou H, Hastie T. Regularization and variable selection via the elastic Net. *J R Stat Soc Ser B Stat Methodol*. 2005;67(2):301–320. doi:10.1111/j.1467-9868.2005.00503.x.
49. Breiman L. Random forests. *Mach Learn*. 2001;45(1):5–32. doi:10.1023/A:1010933404324.
50. Ke G, Meng Q, Finley T, Wang T, Chen W, Ma W, Ye Q, Liu T. LightGBM: a highly efficient gradient boosting decision tree in. In: Ulrike von L, Isabelle G, Samy B, Hanna W, Rob F, editors. *Advances in neural information processing systems*. Long Beach, CA, USA: Curran Associates, Inc; 2017.
51. Cortes C, Vapnik V. Support-vector networks. *Mach Learn*. 1995;20(3):273–297. doi:10.1007/BF00994018.
52. Williams C, Rasmussen C. In: Touretzky D, Mozer MC, Hasselmo M, editors. *Gaussian processes for regression*. In *advances in neural information processing systems*. Cambridge, Massachusetts, USA: MIT Press; 1995. p. 514–520.
53. Kohavi R. A study of cross-validation and bootstrap for accuracy estimation and model selection in. *Proceedings of the 14th International Joint Conference on Artificial Intelligence - ; Vol. Volume 2'95*, IJCAI'Morgan Kaufmann Publishers Inc., 1995. p. 1137–1143.
54. Breiman L. Stacked regressions. *Mach Learn*. 1996;24(1):49–64. doi:10.1007/BF00117832.
55. Application of assisted design of antibody and protein therapeutics (ADAPT) improves efficacy of a Clostridium difficile toxin a single-domain antibody - PubMed. [accessed 2024 May 16]. <https://pubmed.ncbi.nlm.nih.gov/29396522/>.
56. Schymkowitz J, Borg J, Stricher F, Nys R, Rousseau F, Serrano L. The FoldX web server: an online force field. *Nucleic Acids Res*. 2005;33(Web Server):W382–W388. doi:10.1093/nar/gki387.
57. Potapov V, Cohen M, Schreiber G. Assessing computational methods for predicting protein stability upon mutation: good on average but not in the details. *Protein Eng Des Sel*. 2009;22(9):553–560. doi:10.1093/protein/gzp030.
58. Chungyoun M, Ruffolo J, Gray J. Flab: benchmarking deep learning methods for antibody fitness prediction. *Machine Learning for Structural Biology Workshop, NeurIPS 2023*.
59. Ramon A, Ali M, Atkinson M, Saturnino A, Didi K, Visentin C, Ricagno S, Xu X, Greenig M, Sormanni P, et al. Assessing antibody and nanobody nativeness for hit selection and humanization with AbNatiV. *Nat Mach Intell*. 2024;6(1):74–91. doi:10.1038/s42256-023-00778-3.
60. Oeller M, Kang R, Bell R, Ausserwöger H, Sormanni P, Vendruscolo M. Sequence-based prediction of pH-dependent protein solubility using CamSol. *Brief Bioinform*. 2023;24(2):bbad004. doi:10.1093/bib/bbad004.
61. Wolf Pérez AM, Lorenzen N, Vendruscolo M, Sormanni P. Assessment of therapeutic antibody developability by combinations of in vitro and in silico method. *Methods Mol Bio*. 2022;2313:57–113.
62. Hoffmann J, Borgeaud S, Mensch A, Buchatskaya E, Cai T, Rutherford E, de Las Casas D, Hendricks LA, Welbl J, Clark A, et al. Training compute-optimal large language models. [Preprint] 2022. [accessed 2024 May 16]. <http://arxiv.org/abs/2203.15556>.
63. Ying X. An overview of overfitting and its solutions. *J Phys Conf Ser*. 2019;1168:022022. doi:10.1088/1742-6596/1168/2/022022.
64. Hummer AM, Schneider C, Chinery L, Deane CM. Investigating the volume and diversity of data needed for generalizable antibody-antigen $\Delta\Delta G$ prediction. [Preprint] 2023. [accessed 2024 Jul

- 28]. <https://www.biorxiv.org/content/10.1101/2023.05.17.541222v1>.
65. Honegger A, Plu A. Yet another numbering scheme for immunoglobulin variable domains: an automatic modeling and analysis tool. *J Mol Biol.* 2001;309(3):657–670. doi:10.1006/jmbi.2001.4662.
 66. Dunbar J, Deane CM. ANARCI: antigen receptor numbering and receptor classification. *Bioinformatics.* 2016;32(2):298. doi:10.1093/bioinformatics/btv552.
 67. Bateman A, Martin M-J, Orchard S, Magrane M, Ahmad S, Alpi E, Bowler-Barnett EH, Britto R, Bye-A-Jee H, Cukura A. The UniProt consortium, UniProt: the universal protein knowledgebase in 2023. *Nucleic Acids Res.* 2023;51(D1):D523–D531. doi:10.1093/nar/gkac1052.
 68. Olsen TH, Boyles F, Deane CM. Observed antibody space: a diverse database of cleaned, annotated, and translated unpaired and paired antibody sequences. *Protein Sci.* 2022;31(1):141–146. doi:10.1002/pro.4205.
 69. Indi—integrated nanobody database for immunoinformatics | nucleic acids research | oxford academic. [accessed 2024 May 8]. <https://academic.oup.com/nar/article/50/D1/D1273/6423188>.
 70. Dunbar J, Krawczyk K, Leem J, Baker T, Fuchs A, Georges G, Shi J, Deane CM. SAbDab: the structural antibody database. *Nucleic Acids Res.* 2014;42(D1):D1140–D1146. doi:10.1093/nar/gkt1043.
 71. Pedregosa F, Varoquaux G, Gramfort A, Michel V, Thirion B, Grisel O, Blondel M, Müller A, Nothman J, Louppe G, Prettenhofer P, et al. Scikit-learn: machine learning in Python. [Preprint] (2018). [accessed 2024 May 9]. <http://arxiv.org/abs/1201.0490>.
 72. Protein Data Bank | Nucleic Acids Research | Oxford Academic. [accessed 2024 May 9]. <https://academic.oup.com/nar/article/28/1/235/2384399?login=false>.
 73. Santoro MM, Bolen DW. Unfolding free energy changes determined by the linear extrapolation method. I. Unfolding of phenylmethanesulfonyl alpha-chymotrypsin using different denaturants. *Biochemistry.* 1988;27(21):8063–8068. doi:10.2307/3948866.
 74. Satopaa V, Albrecht J, Irwin D, Raghavan B. Finding a “kneedle” in a haystack: detecting knee points in system behavior in 2011. 31st International conference on distributed computing Systems workshops. 2011. p. 166–171.
 75. Alvarez JAE, Dean SN. TEMPRO: nanobody melting temperature estimation model using protein embeddings. *Sci Rep.* 2024;14(1):19074. doi:10.1038/s41598-024-70101-6.
 76. Gómez-Mulas A, Naganathan AN, Pey AL. The lack of trade-off between conformational stability and binding affinity in a nanobody with therapeutic potential for a misfolding disease. [Preprint]. 2024. [accessed 2024 Nov 26]. <https://www.biorxiv.org/content/10.1101/2024.09.15.612864v1>.
 77. Gómez-Mulas A, Salido E, Pey AL. Thermodynamic versus kinetic basis for the high conformational stability of nanobodies for therapeutic applications. [Preprint] (2024). Accessed 2024 Nov 26]. <https://www.biorxiv.org/content/10.1101/2024.08.28.609651v1>.