

Partitioning, representation, and automation in canonical and non-canonical peptide modelling

Speaker: **Raúl Fernández-Díaz (PhD Candidate UCD – IBM Research)**

UCD: R. Cossio-Pérez, C. Agoni, D.C. Shields

IBM Research: T.L. Hoang, V. Lopez

Novo Nordisk: R. Ochoa

Part 0 - Introduction



More information
and contact info



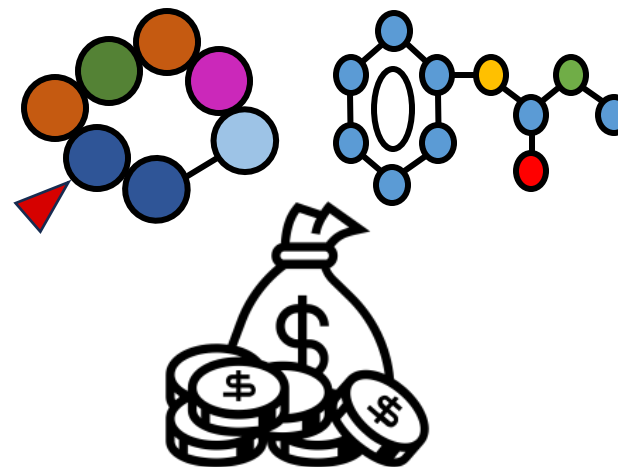
A tale of two peptides

Natural (or canonical) peptides Synthetic peptides (or non-canonical) and peptidomimetics

- Protein sequences (20 aa)
- Cheap
- Not great drugs

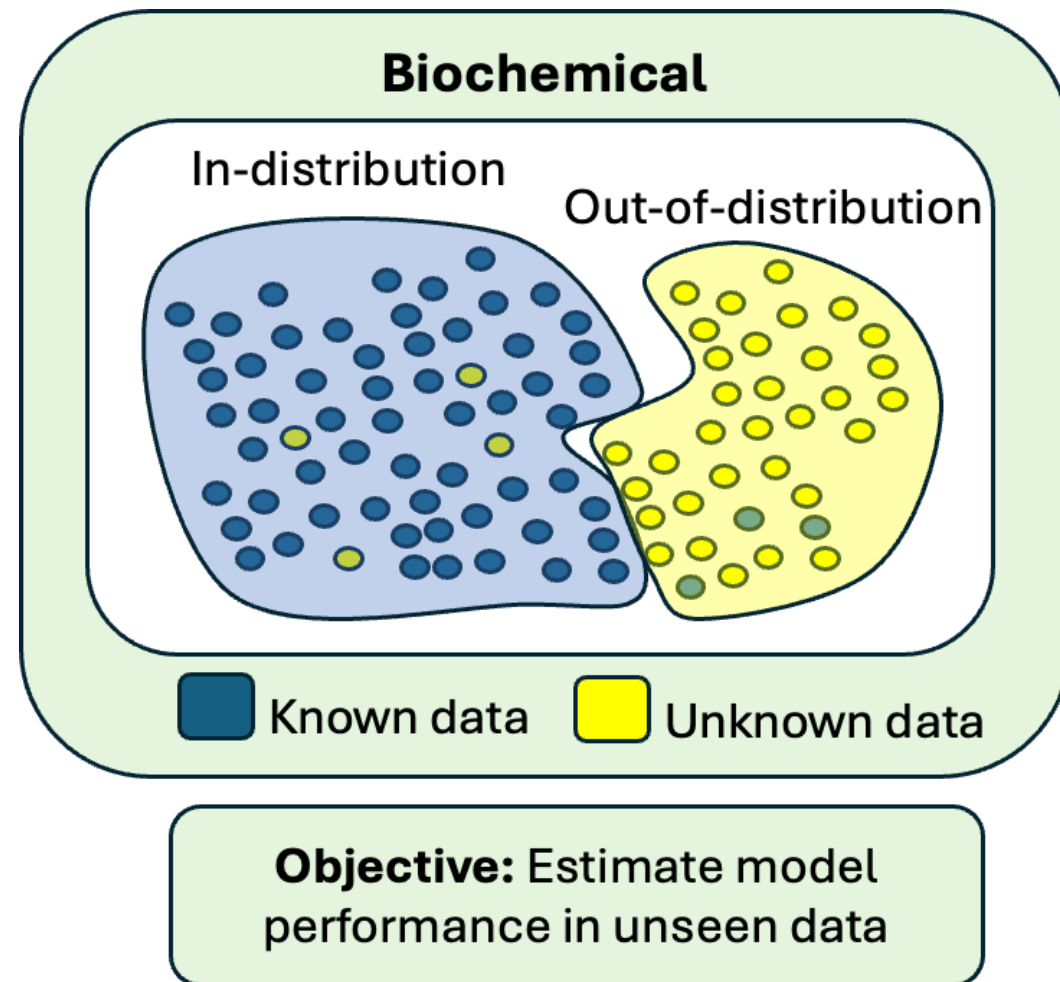
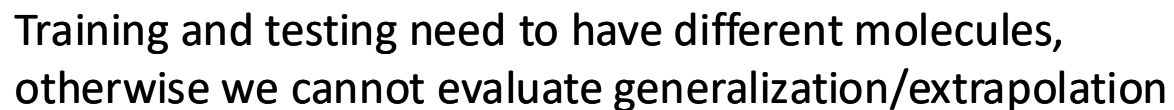


- Modified peptides + small molecules
- Expensive
- Better drug candidates





Central Assumption of ML: “Training data is representative of prediction data”

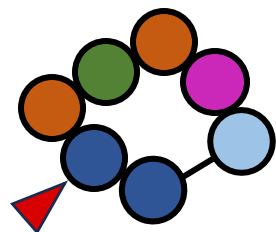
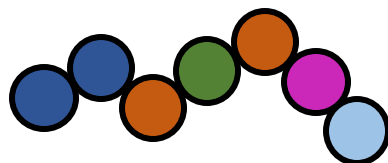


More information
and contact info



Peptide representation

We need to transform molecules into numerical vectors



Methods



[1, 0.3, 0.4, 2.1, ...]

- **Heuristic:** Molecular FPs
 - Substructure-based: ECFP/Avalon
 - Monomer sequence: PepFuNN
- **Data-driven:** Pre-trained representation models
 - [Protein/Chemical/Peptide] Language Models
 - GNNs

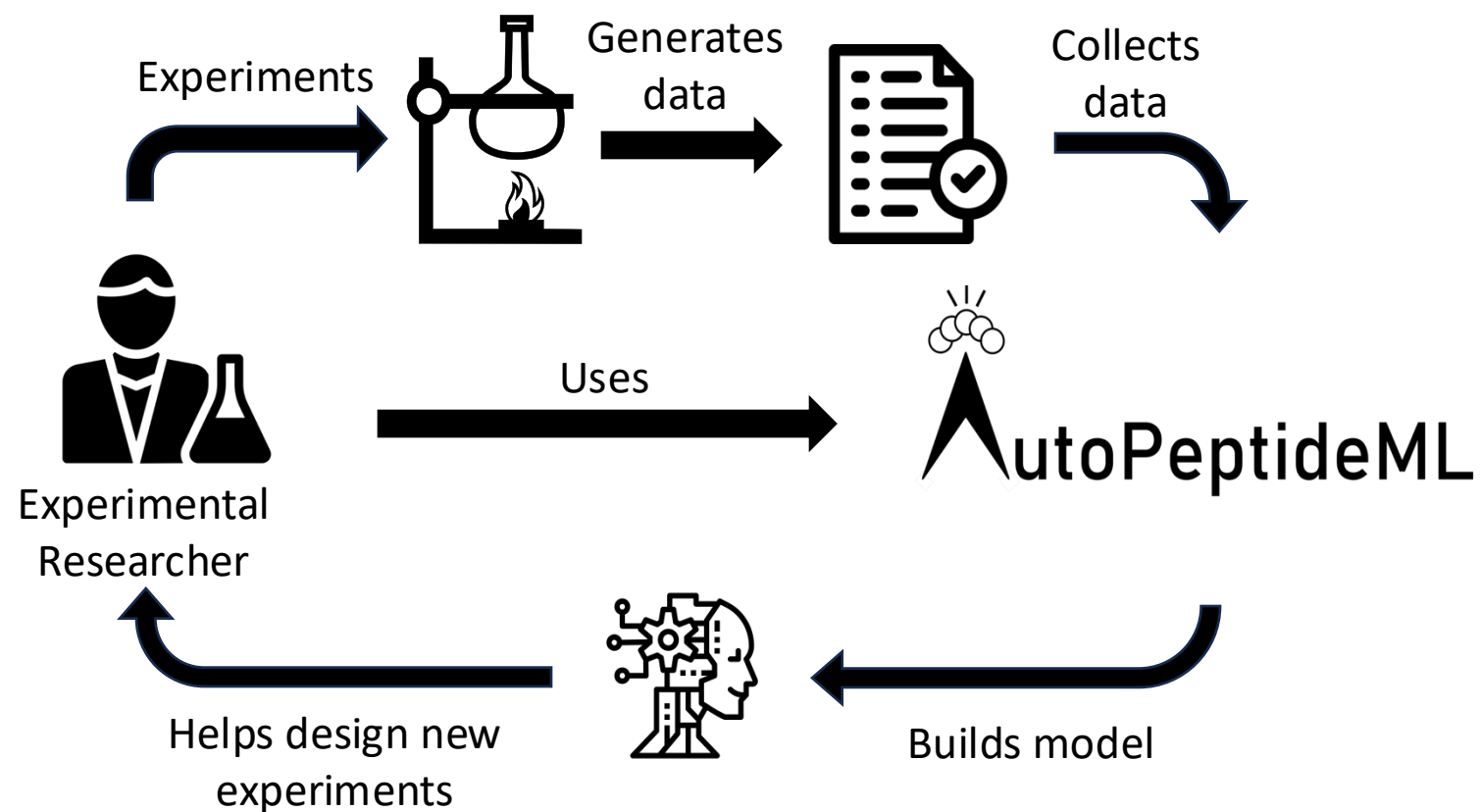
More information
and contact info



Automation



R. Fernández-Díaz et al.,
AutoPeptideML: a study on how
to build more trustworthy
peptide bioactivity
predictors, *Bioinformatics*,
Volume 40, Issue 9, September
2024, btae555



Design Requirements

1. Robust evaluation
2. Reproducibility
3. Easier to integrate with experimental workflow



Objectives

1. How to automatically build peptide property prediction models (and how to evaluate them)
2. How to extrapolate from natural to synthetic peptides or peptidomimetics

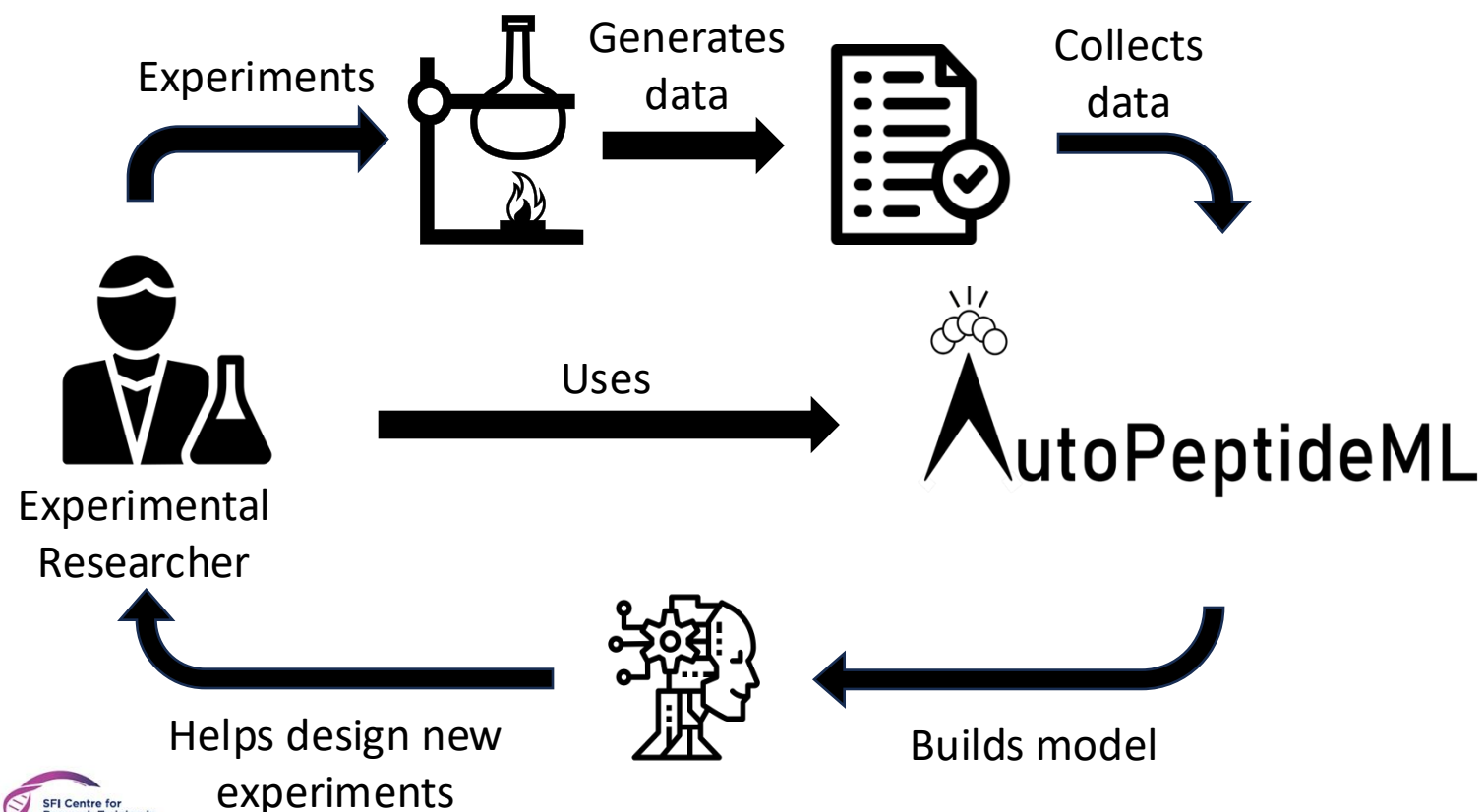


Part 1 - AutoPeptideML

The background of the slide is a gradient of purple and blue. On the right side, there is a complex, abstract graphic consisting of numerous small, dark purple and blue dots (nodes) connected by thin, light purple lines (edges). These nodes and lines are arranged in a way that suggests a network or a molecular structure. Overlaid on this network are several large, flowing, wavy lines in shades of purple and blue, creating a sense of movement and depth. The overall aesthetic is modern and scientific.



Automation



Design Requirements

1. Competitive performance
2. Reliable evaluation so that experimental scientist can trust the models
3. Easy to use

More information
and contact info



Collected **18**
datasets used for
building different
peptide bioactivity
predictors

How good can automation really be?



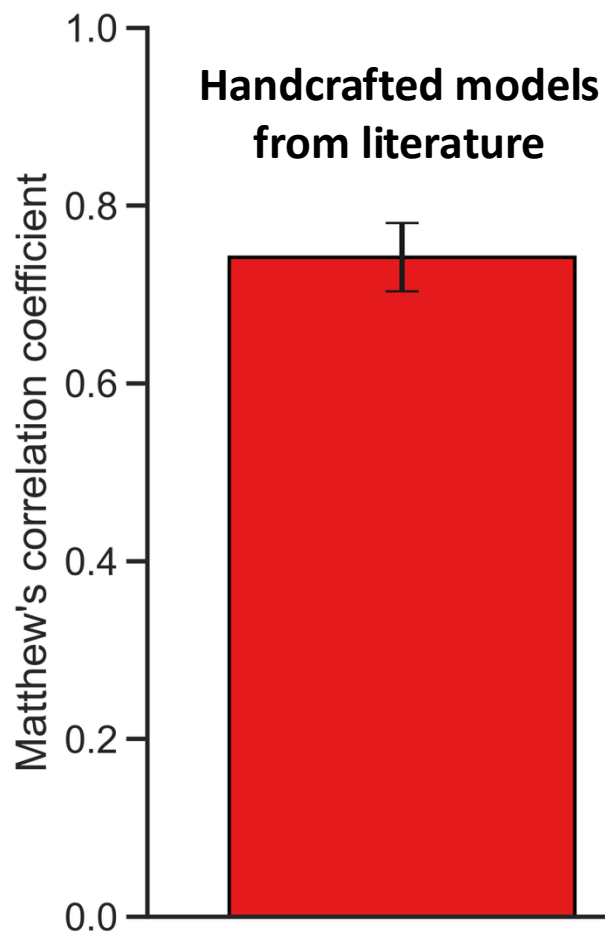
R. Fernández-Díaz et al.,
AutoPeptideML: a study on how
to build more trustworthy
peptide bioactivity
predictors, *Bioinformatics*,
Volume 40, Issue 9, September
2024, btae555

More information
and contact info



Collected **18**
datasets used for
building different
peptide bioactivity
predictors

How good can automation really be?



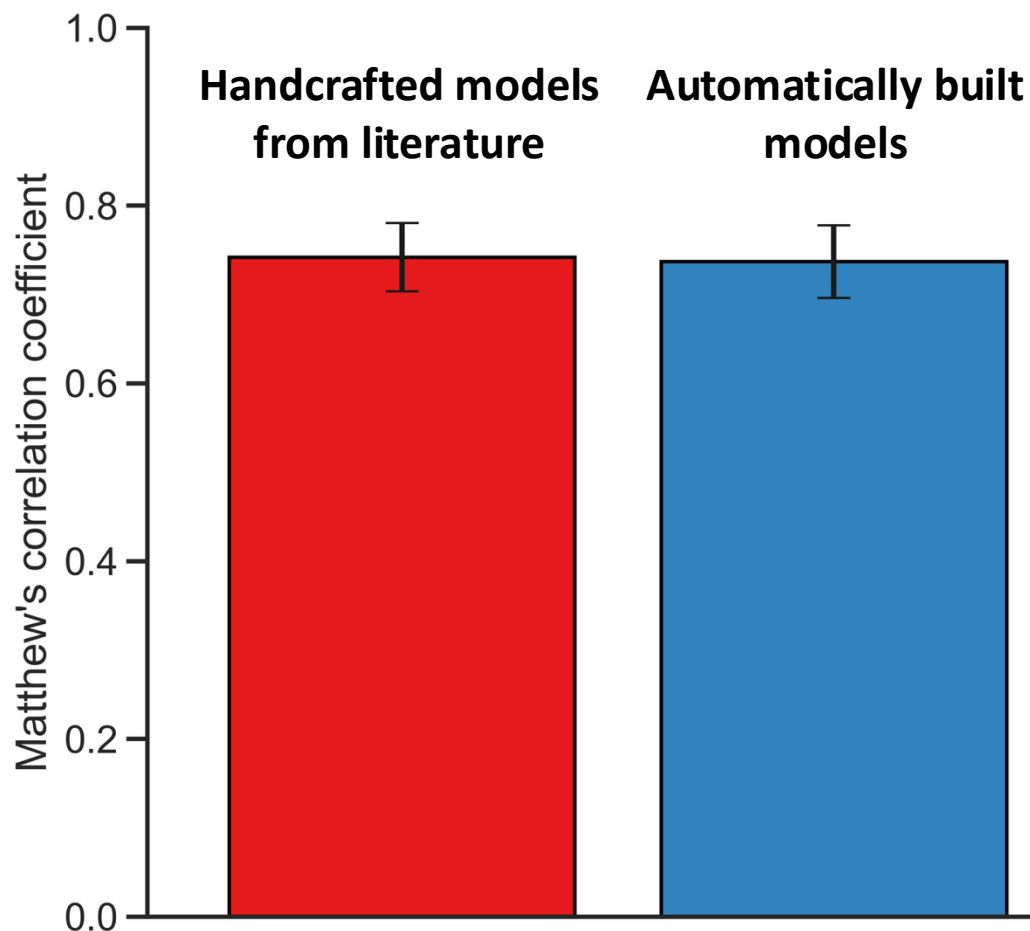
R. Fernández-Díaz et al.,
AutoPeptideML: a study on how
to build more trustworthy
peptide bioactivity
predictors, *Bioinformatics*,
Volume 40, Issue 9, September
2024, btae555

More information
and contact info



Collected **18**
datasets used for
building different
peptide bioactivity
predictors

Automation is as good as human developers



Low intensity computing



OPTUNA

Bayesian Optimization for
hyperparameter selection

**Protein Language
Models**

General representation/
featurization



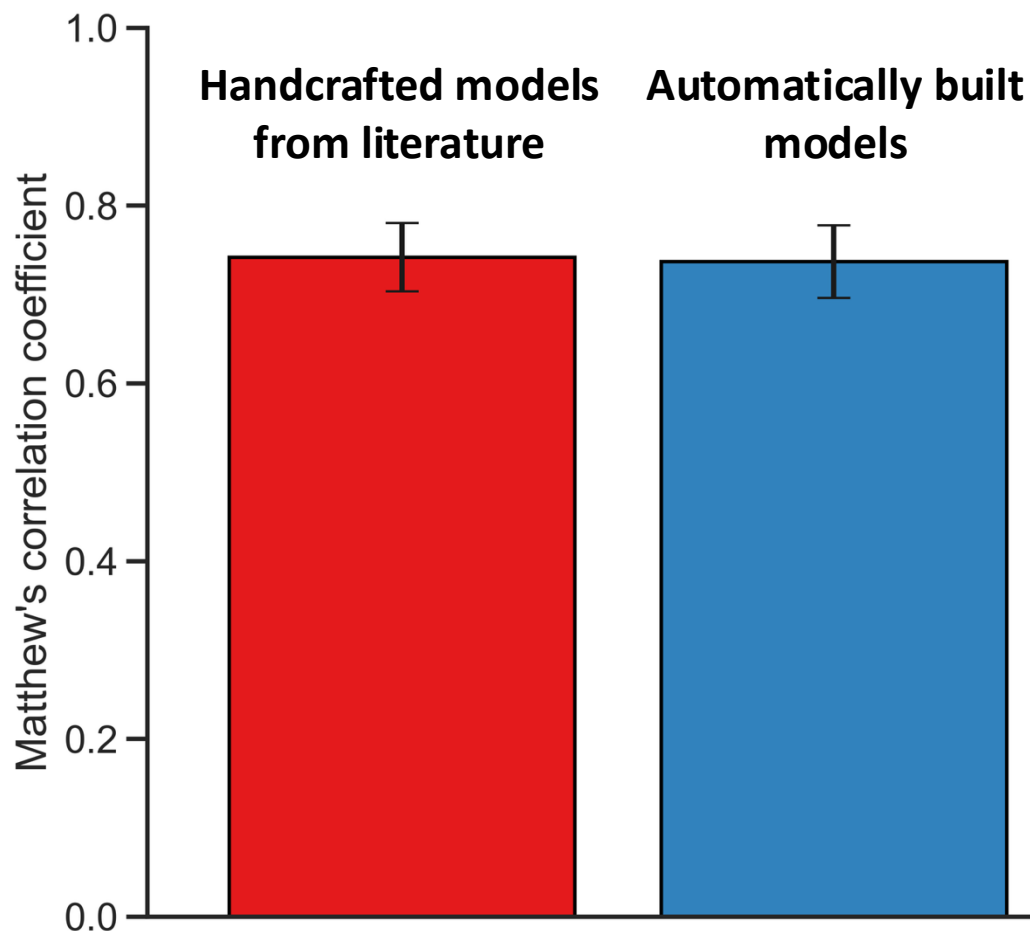
R. Fernández-Díaz et al.,
AutoPeptideML: a study on how
to build more trustworthy
peptide bioactivity
predictors, *Bioinformatics*,
Volume 40, Issue 9, September
2024, btae555

More information
and contact info



Collected **18**
datasets used for
building different
peptide bioactivity
predictors

How important is the choice of negative samples?



Before:

- Random peptides
- Peptides from Uniprot
- Protein fragments
- Scrambled sequences

Problem: They might differ from positives due to confounding biophysical characteristics (e.g. solubility, membrane crossing, etc.)

Now:

- Other bioactive peptides

Peptipedia



R. Fernández-Díaz et al.,
AutoPeptideML: a study on how
to build more trustworthy
peptide bioactivity
predictors, *Bioinformatics*,
Volume 40, Issue 9, September
2024, btae555

More information
and contact info



Collected **18**
datasets used for
building different
peptide bioactivity
predictors

Proper choice of negatives is quite important

Before:

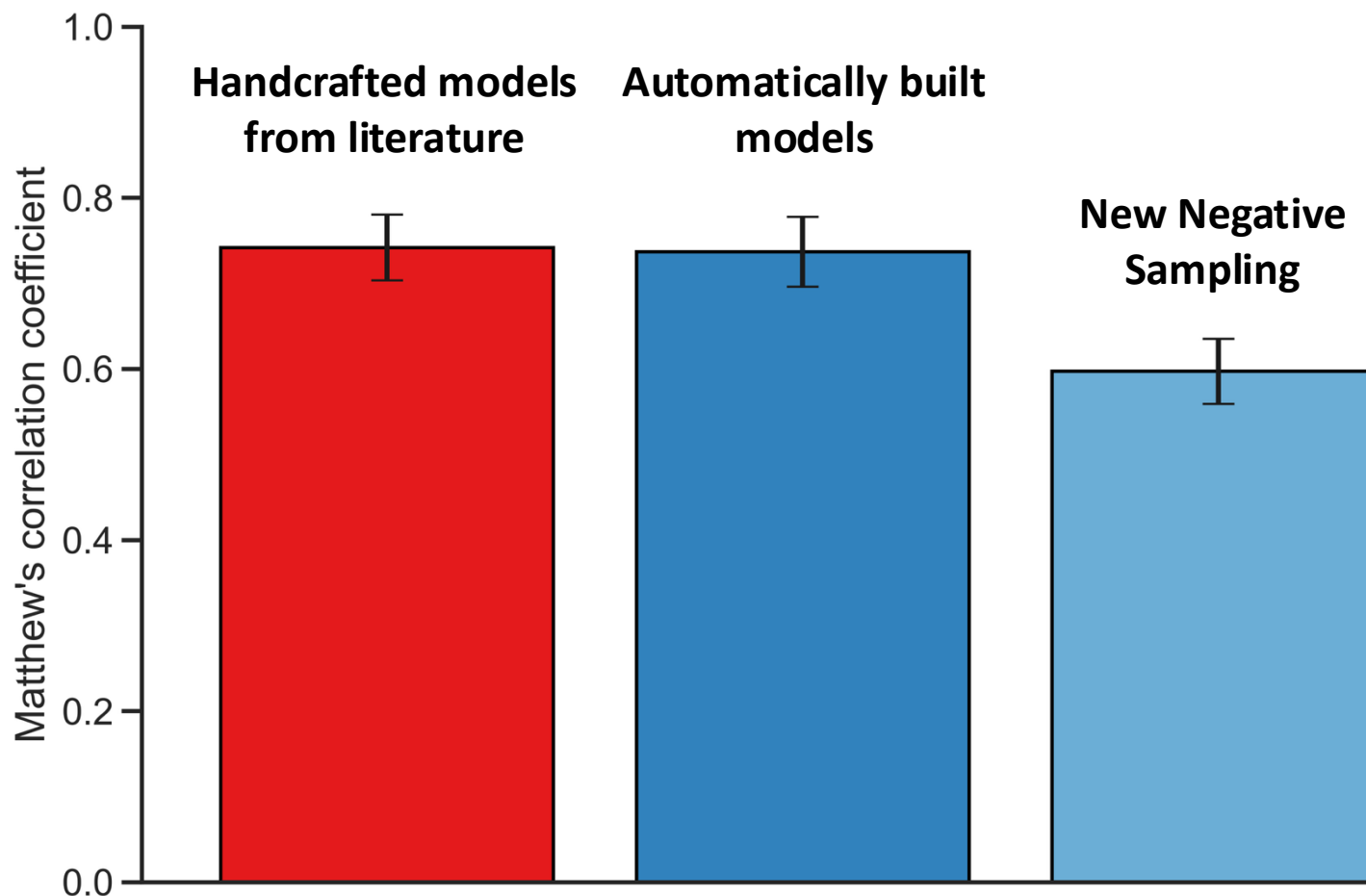
- Random peptides
- Peptides from Uniprot
- Protein fragments
- Scrambled sequences

Problem: They might differ from positives due to confounding biophysical characteristics (e.g. solubility, membrane crossing, etc.)

Now:

- Other bioactive peptides

Peptipedia



 **AutoPeptideML**

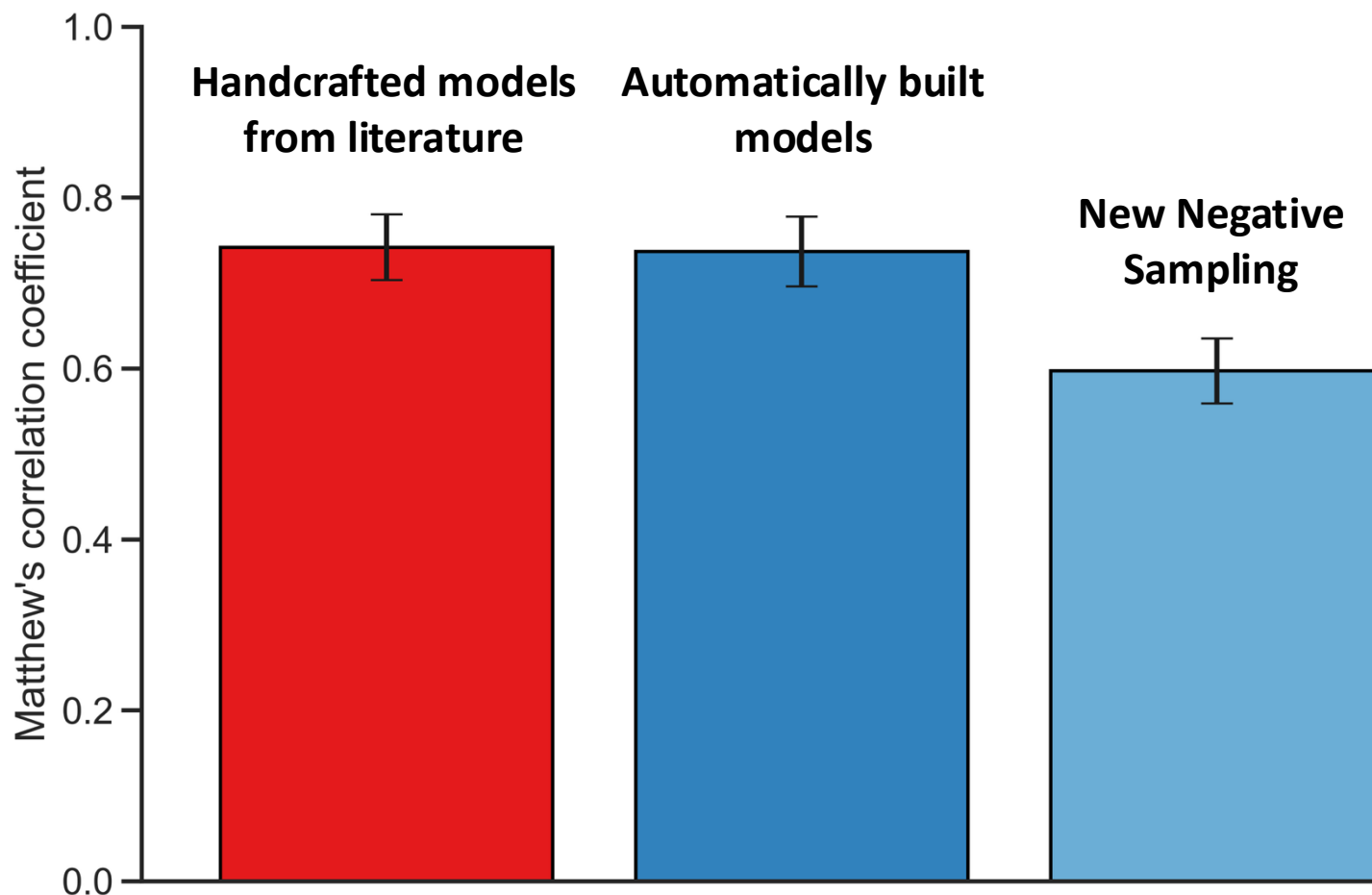
R. Fernández-Díaz et al.,
AutoPeptideML: a study on how
to build more trustworthy
peptide bioactivity
predictors, *Bioinformatics*,
Volume 40, Issue 9, September
2024, btae555

More information
and contact info

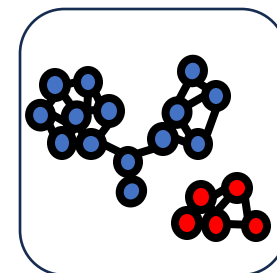


How important is to have independent test sets?

Collected **18**
datasets used for
building different
peptide bioactivity
predictors



Testing data
should only
include unseen
molecules not
similar to
training
(sequence
alignment)



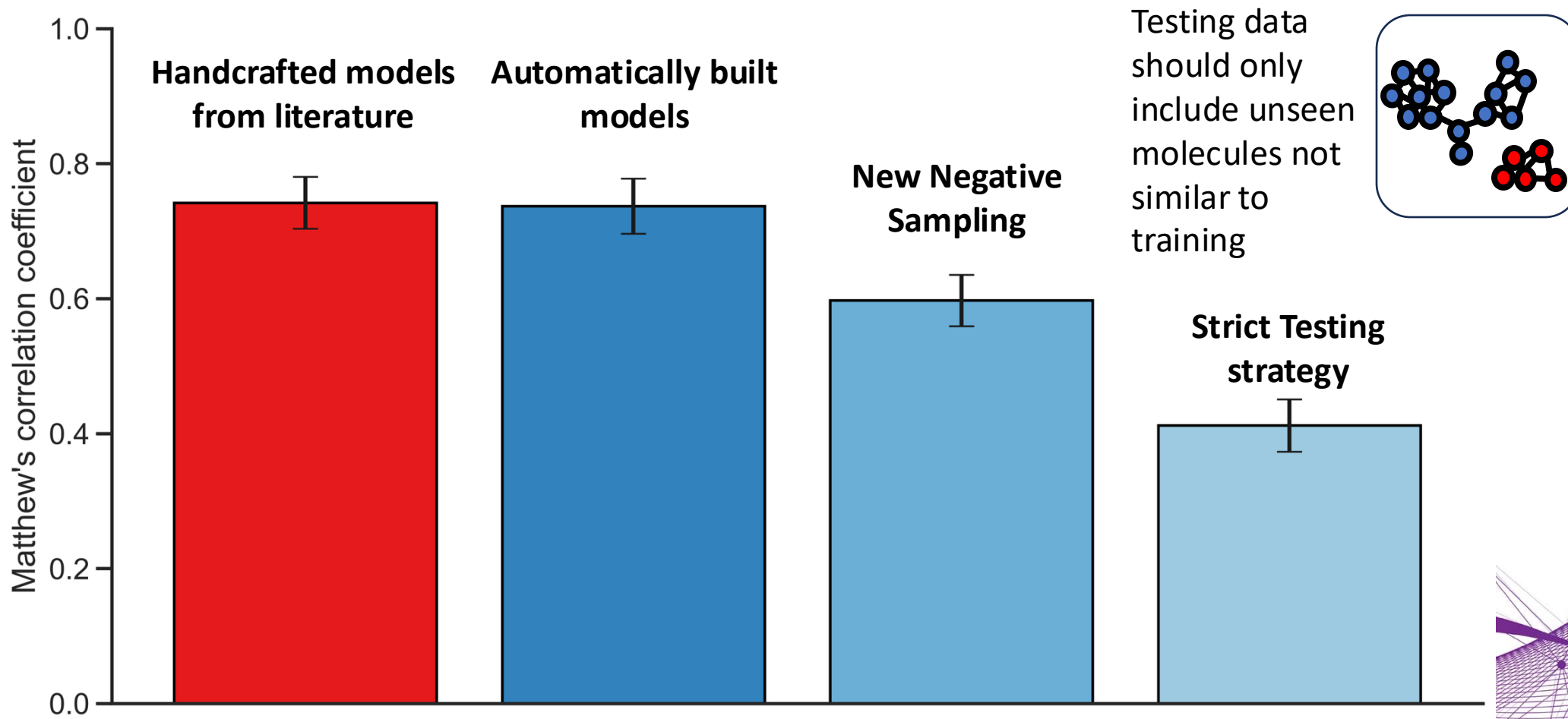
R. Fernández-Díaz et al.,
AutoPeptideML: a study on how
to build more trustworthy
peptide bioactivity
predictors, *Bioinformatics*,
Volume 40, Issue 9, September
2024, btae555

More information
and contact info



Collected **18**
datasets used for
building different
peptide bioactivity
predictors

Independence in test sets is crucial



 AutoPeptideML

R. Fernández-Díaz et al.,
AutoPeptideML: a study on how
to build more trustworthy
peptide bioactivity
predictors, *Bioinformatics*,
Volume 40, Issue 9, September
2024, btae555

More information
and contact info



Automating ML for natural peptides



R. Fernández-Díaz et al.,
AutoPeptideML: a study on how
to build more trustworthy
peptide bioactivity
predictors, *Bioinformatics*,
Volume 40, Issue 9, September
2024, btae555

Webserver - GUI



BIOINFORMATICS PUBLICATION CHEMRXIV PREPRINT GITHUB REPOSITORY
DOCUMENTATION

Welcome to the AutoPeptideML webserver.

The next steps will help you build your own model.

0. Modelling task

First, start by defining the prediction task.

What is the prediction problem you are facing?

Classification (categorical values)

1. Inputs

In this section you will define the data from which you want the model to learn.

☐ Download sample dataset

Please upload dataset with your peptides and their labels if available

Drag and drop file here
Limit 200MB per file

Browse files

CLI tool

```
| AutoPeptideML v.2.0.3 |  
| By Raul Fernandez-Diaz |
```

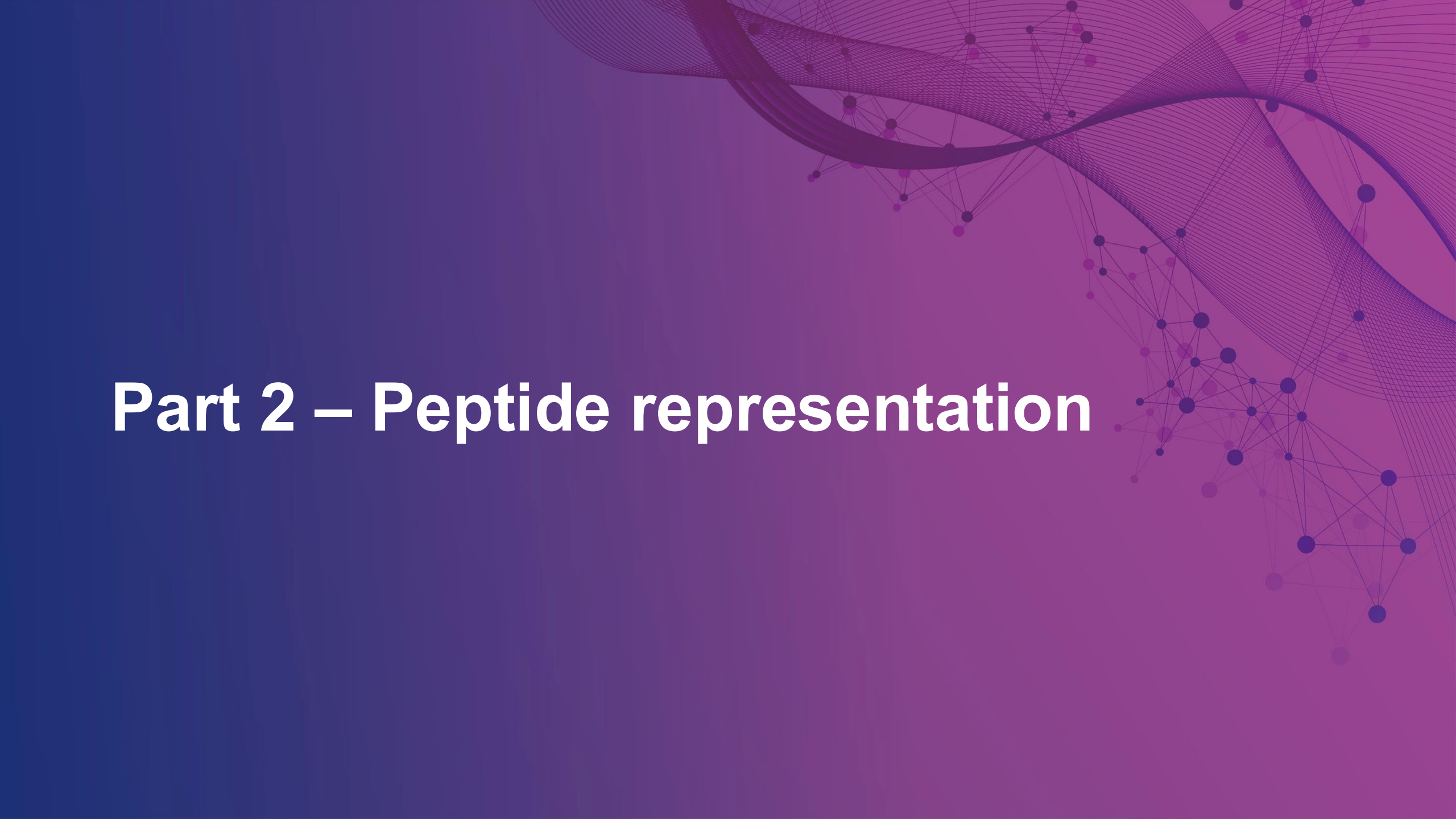
Model builder

```
Part 1 - Define the data and preprocessing steps  
[?] What is the modelling problem you're facing?:  
> Classification (returning categorical value)  
Regression(returnin continuous value)
```

Python Package

```
df = pd.read_csv(osp.join(PATH, 'original_data', f'c-{dataset}.csv'))  
apml = AutoPeptideML(  
    data=df,  
    outputdir=f'apml-{dataset}',  
    sequence_field='SMILES',  
    label_field='labels'  
)  
apml.build_models(  
    task='class',  
    reps=['esm2-8m', 'peptideclm', 'chemberta-2', 'ecfp-16'],  
    models=['svm', 'knn', 'rf', 'lightgbm', 'xgboost'],  
    device='mps',  
    n_trials=10  
)  
apml.create_report()  
return apml
```

Part 2 – Peptide representation

The background of the slide features a complex, abstract design. It consists of several thick, wavy, translucent lines in shades of purple and blue that flow across the frame. Overlaid on these lines is a network of small, dark-colored dots. These dots are interconnected by thin, light-colored lines, creating a web-like structure that resembles a molecular or data network. The overall aesthetic is modern and scientific, with a color palette dominated by deep purples and blues.



Can we leverage data on cheaper experiments to prioritise more expensive experiments?



Train



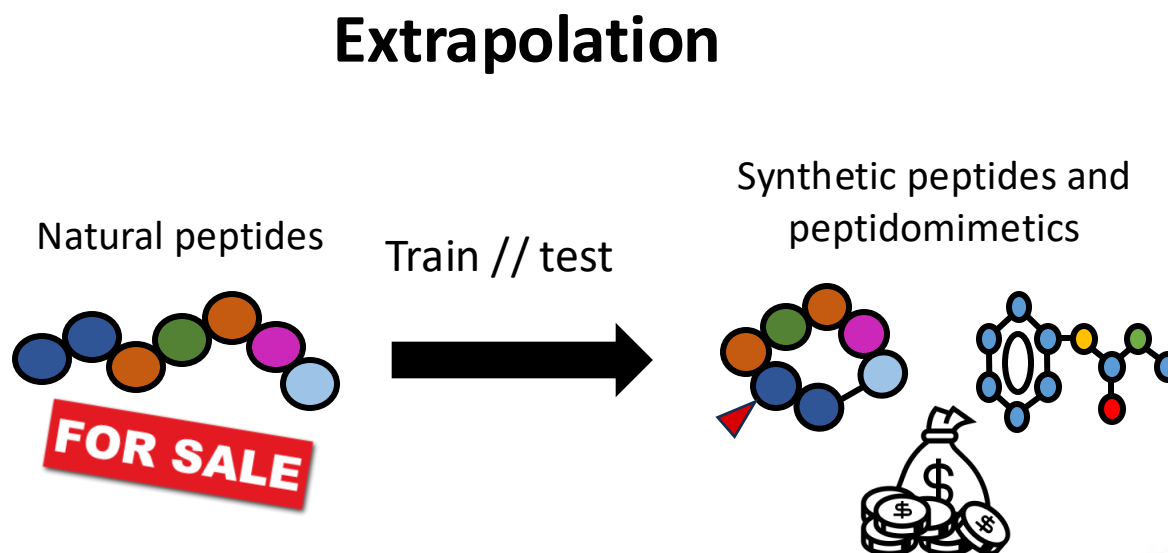
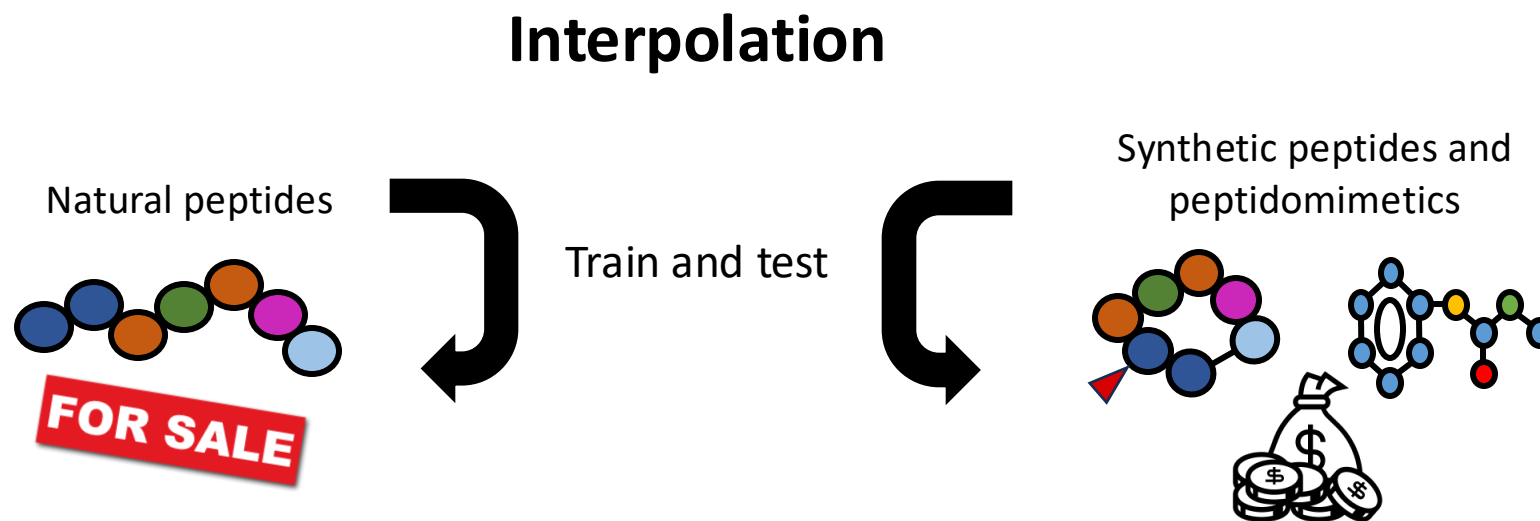
Predict

Synthetic peptides and peptidomimetics
(expensive but better drugs)





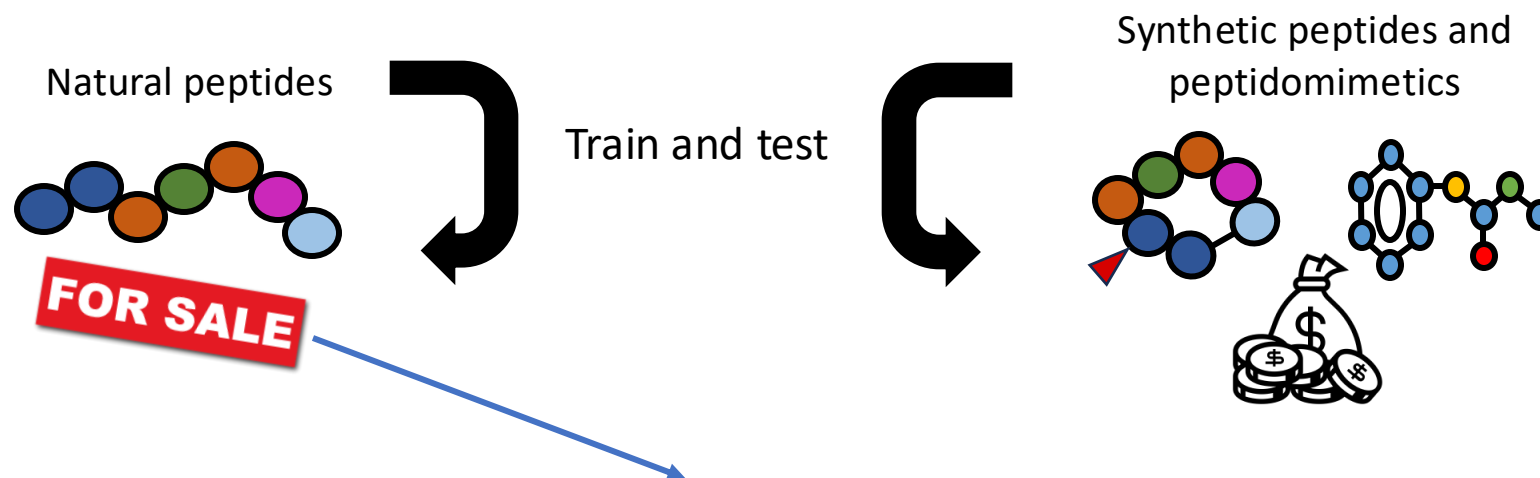
Computational experiments





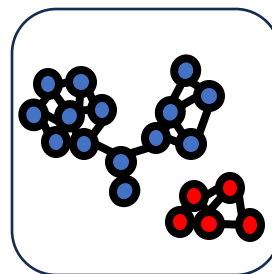
Computational experiments

Interpolation



How to build train/test splits: Should we use sequence alignment for measuring peptide similarity?

Similarity-informed
train/test split





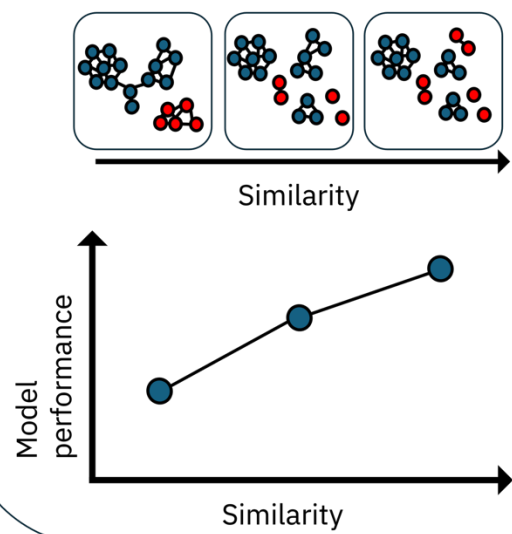
Finding the best similarity metric

Fernandez-Diaz R, et al. A new framework for evaluating model out-of-distribution generalisation for the biochemical domain. In The Thirteenth International Conference on Learning Representations 2025.

Hestia-GOOD framework

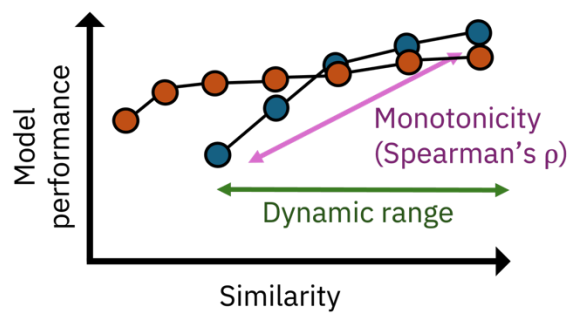
GOOD curve

Model performance as a function of train-test similarity



Similarity metric selection

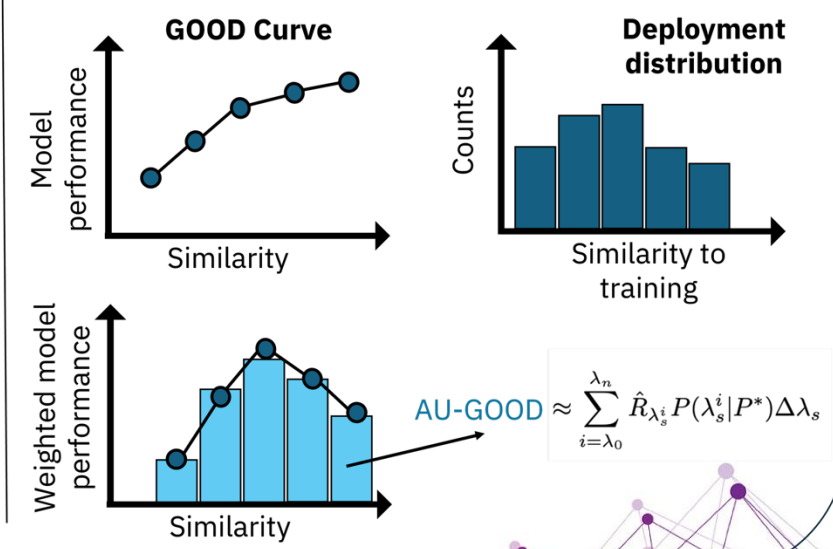
Quantitative analysis of best similarity function for a given task/dataset



- **Monotonicity:** Is model performance a function of train-test similarity?
- **Dynamic range:** What is the resolution of the similarity metric?

AU-GOOD metric

Estimation of model performance conditioned on a deployment distribution





Best metrics for each dataset

Metrics explored:

- ECFP various radii
- MAPc various radii
- MMSeqs2 (alignment)
- Needleman-Wunsch (alignment)
- ESM2-8M embedding distance
- Molformer-XL embedding distance

- Analysis of 8 datasets:
 - 4 natural
 - 4 synthetic
- Chemical FPs are better than sequence alignment for natural peptides





Best metrics for each dataset

Metrics explored:

- ECFP various radii
- MAPc various radii
- MMSeqs2 (alignment)
- Needleman-Wunsch (alignment)
- ESM2-8M embedding distance
- Molformer-XL embedding distance

- Analysis of 8 datasets:
 - 4 natural
 - 4 synthetic
- **Chemical FPs are better than sequence alignment for natural peptides**

Dataset	Peptide type	Task	Similarity Type	Similarity	Dynamic range (↑) [a]	Monotonicity (↑) [b]
Protein-peptide binding affinity	Standard	Regression	Chemical FP	MAPc-8	70 %	0.8 ± 0.1
Protein-peptide binding affinity	Modified	Regression	Chemical FP	MAPc-20	80 %	0.95 ± 0.03
Cell penetration	Standard	Classification	Chemical FP	MAPc-8	60 %	0.98 ± 0.04
Cell penetration	Modified	Classification	Chemical FP	MAPc-12	60 %	0.5 ± 0.2
Antibacterial	Standard	Classification	Chemical FP	MAPc-8	60 %	0.97 ± 0.02
Antibacterial	Modified	Classification	Chemical FP	ECFP-12	50 %	0.9 ± 0.1
Antiviral	Standard	Classification	Sequence Alignment	MMSeqs2	80 %	0.96 ± 0.05
Antiviral	Modified	Classification	Chemical FP	MAPc-12	70 %	0.6 ± 0.2



More information
and contact info



Fernandez-Diaz R, et al. A new framework for evaluating model out-of-distribution generalisation for the biochemical domain. InThe Thirteenth International Conference on Learning Representations 2025.

Best metrics for each dataset

Metrics explored:

- ECFP various radii
 - MAPc various radii
 - MMSeqs2 (alignment)
- Needleman-Wunsch (alignment)
 - ESM2-8M embedding distance
 - Molformer-XL embedding distance

- Analysis of 8 datasets:
 - 4 natural
 - 4 synthetic
- **Chemical FPs are better than sequence alignment for natural peptides**

Dataset	Peptide type	Task	Similarity Type	Similarity	Dynamic range (↑) [a]	Monotonicity (↑) [b]
Protein-peptide binding affinity	Standard	Regression	Chemical FP	MAPc-8	70 %	0.8 ± 0.1
Protein-peptide binding affinity	Modified	Regression	Chemical FP	MAPc-20	80 %	0.95 ± 0.03
Cell penetration	Standard	Classification	Chemical FP	MAPc-8	60 %	0.98 ± 0.04
Cell penetration	Modified	Classification	Chemical FP	MAPc-12	60 %	0.5 ± 0.2
Antibacterial	Standard	Classification	Chemical FP	MAPc-8	60 %	0.97 ± 0.02
Antibacterial	Modified	Classification	Chemical FP	ECFP-12	50 %	0.9 ± 0.1
Antiviral	Standard	Classification	Sequence Alignment	MMSeqs2	80 %	0.96 ± 0.05
Antiviral	Modified	Classification	Chemical FP	MAPc-12	70 %	0.6 ± 0.2



More information
and contact info


Fernandez-Diaz R, et al. A new framework for evaluating model out-of-distribution generalisation for the biochemical domain. InThe Thirteenth International Conference on Learning Representations 2025.

Best metrics for each dataset

Metrics explored:

- ECFP various radii
 - MAPc various radii
 - MMSeqs2 (alignment)

- Needleman-Wunsch (alignment)
 - ESM2-8M embedding distance
 - Molformer-XL embedding distance

- Analysis of 8 datasets:
 - 4 natural
 - 4 synthetic
 - Chemical FPs are better than sequence alignment for natural peptides

Dataset	Peptide type	Task	Similarity Type	Similarity	Dynamic range (↑) [a]	Monotonicity (↑) [b]
Protein-peptide binding affinity	Standard	Regression	Chemical FP	MAPc-8	70 %	0.8 ± 0.1
Protein-peptide binding affinity	Modified	Regression	Chemical FP	MAPc-20	80 %	0.95 ± 0.03
Cell penetration	Standard	Classification	Chemical FP	MAPc-8	60 %	0.98 ± 0.04
Cell penetration	Modified	Classification	Chemical FP	MAPc-12	60 %	0.5 ± 0.2
Antibacterial	Standard	Classification	Chemical FP	MAPc-8	60 %	0.97 ± 0.02
Antibacterial	Modified	Classification	Chemical FP	ECFP-12	50 %	0.9 ± 0.1
Antiviral	Standard	Classification	Sequence Alignment	MMSeqs2	80 %	0.96 ± 0.05
Antiviral	Modified	Classification	Chemical FP	MAPc-12	70 %	0.6 ± 0.2

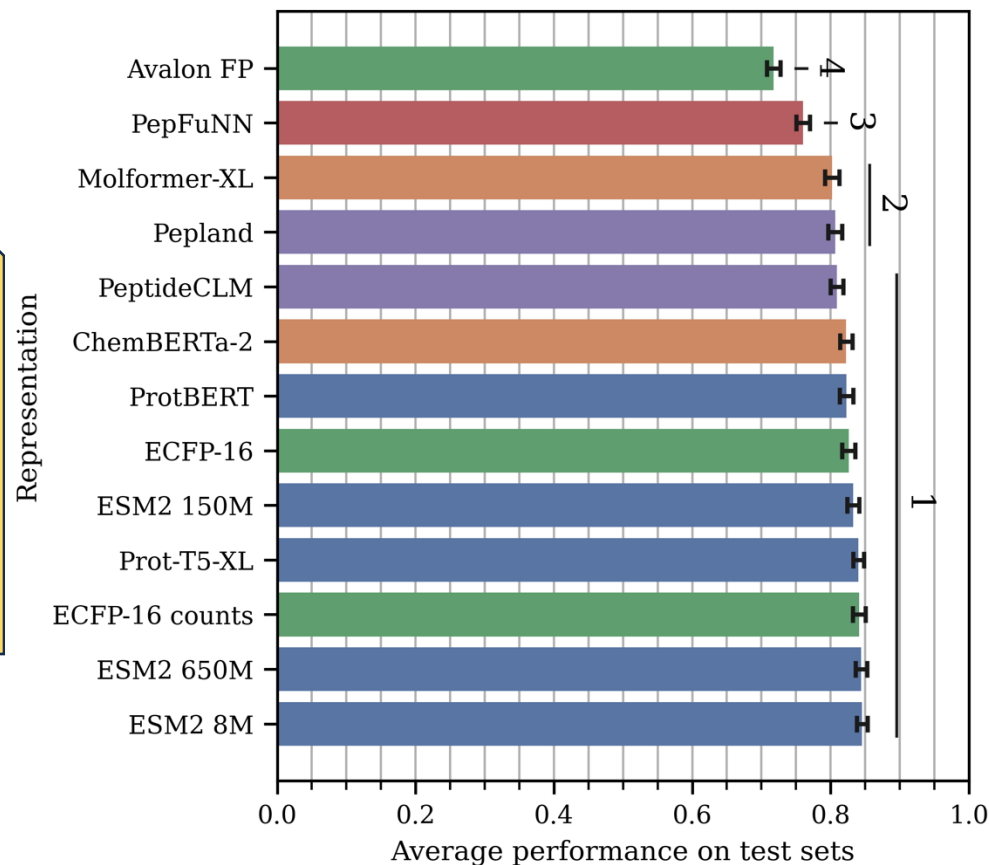




Interpolation experiments

Synthetic to synthetic

Natural to natural



- Average on 4 datasets.
- Statistical analysis are Kruskal-Wallis with *post-hoc* Wilcoxon test.
- Significance defined with Bonferroni correction.



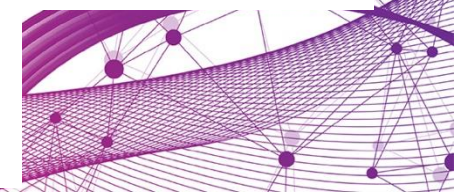
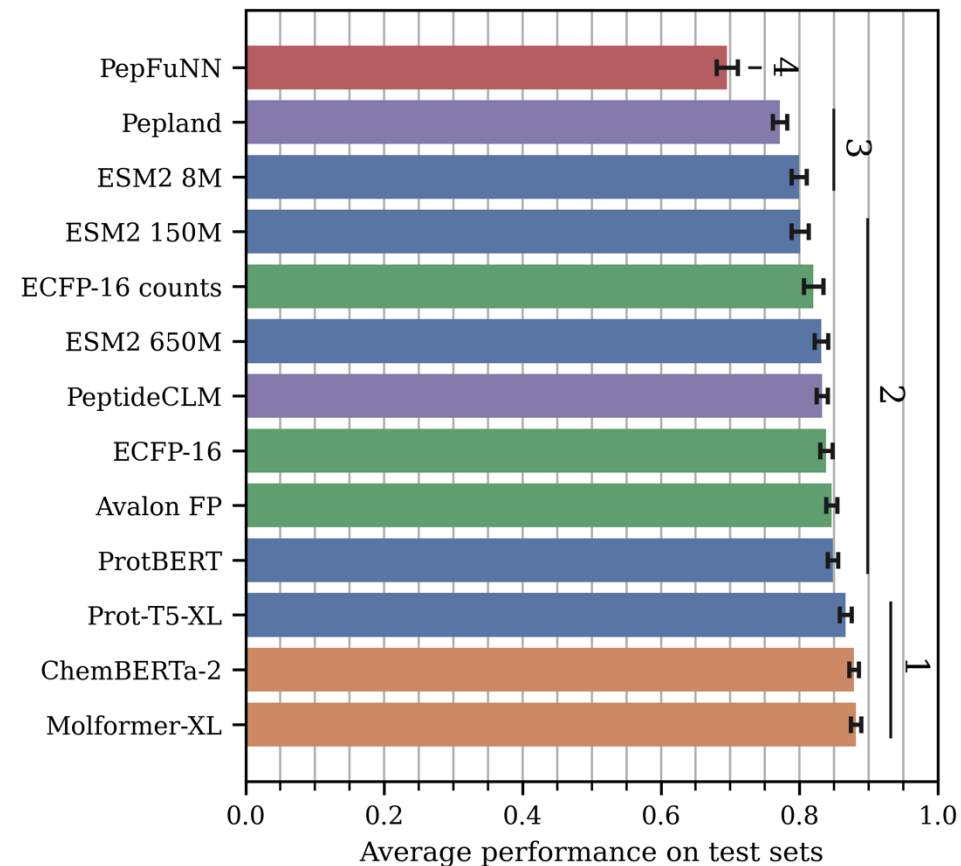


Interpolation experiments

Natural to natural

Synthetic to synthetic

- Average on 4 datasets.
- Statistical analysis are Kruskal-Wallis with *post-hoc* Wilcoxon test.
- Significance defined with Bonferroni correction.

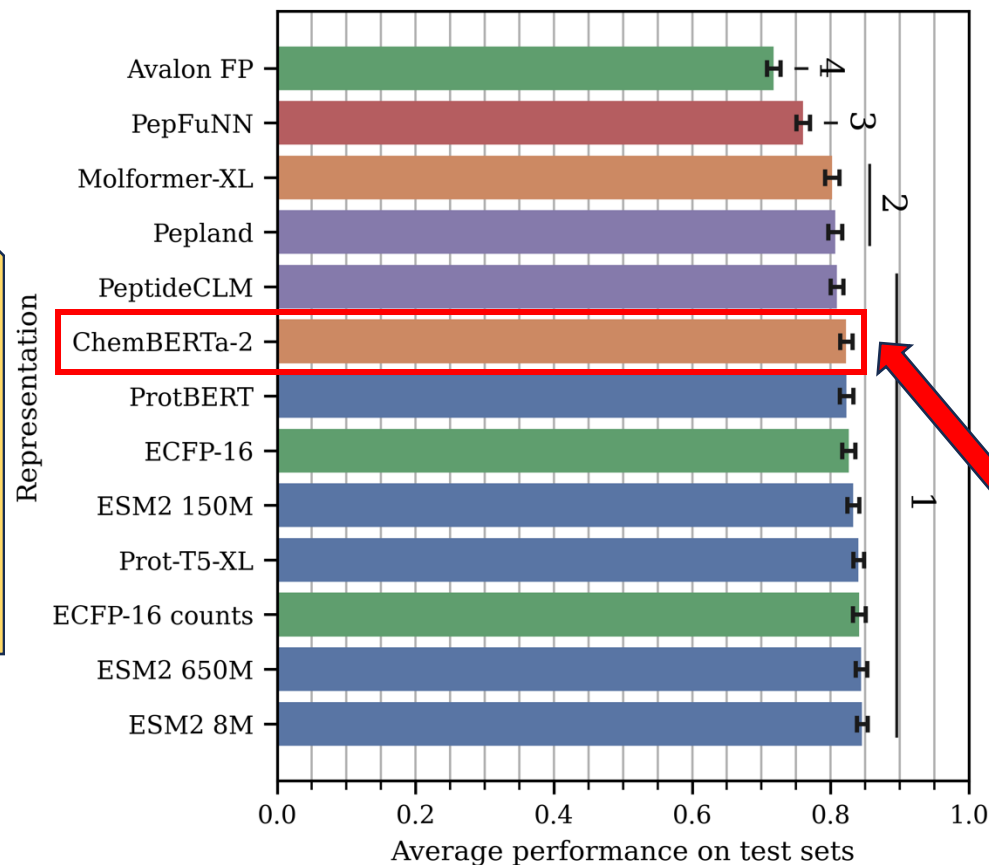




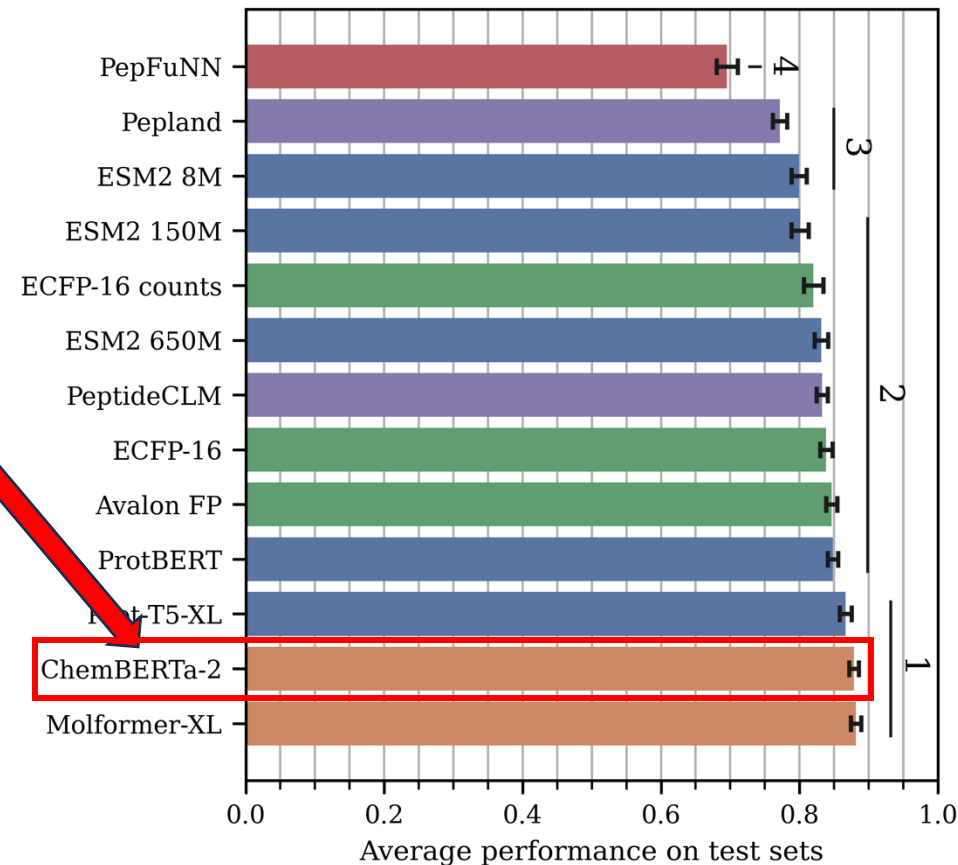
Interpolation experiments

- Average on 4 datasets.
- Statistical analysis are Kruskal-Wallis with *post-hoc* Wilcoxon test.
- Significance defined with Bonferroni correction.

Natural to natural



Synthetic to synthetic



Representation Family

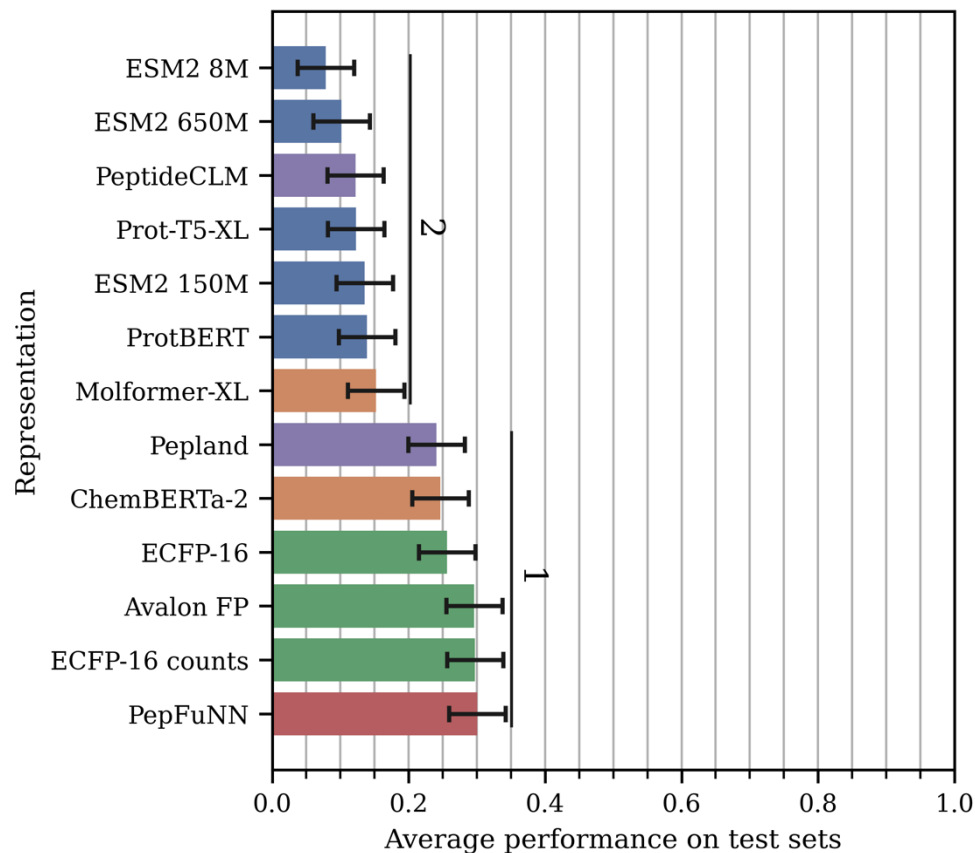
PLM CLM Chemical FP Peptide FP Peptide LM/GNN





Natural to synthetic extrapolation

- Average on 4 datasets.
- Statistical analysis are Kruskal-Wallis with *post-hoc* Wilcoxon test.
- Significance defined with Bonferroni correction.



Representation Family

PLM CLM Chemical FP Peptide FP Peptide LM/GNN

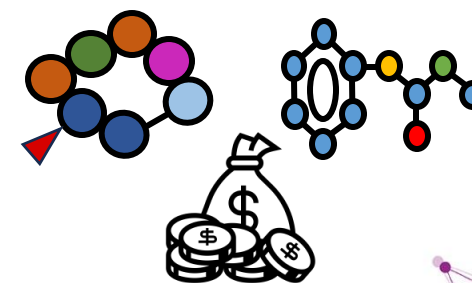
Natural peptides



FOR SALE

Train // test

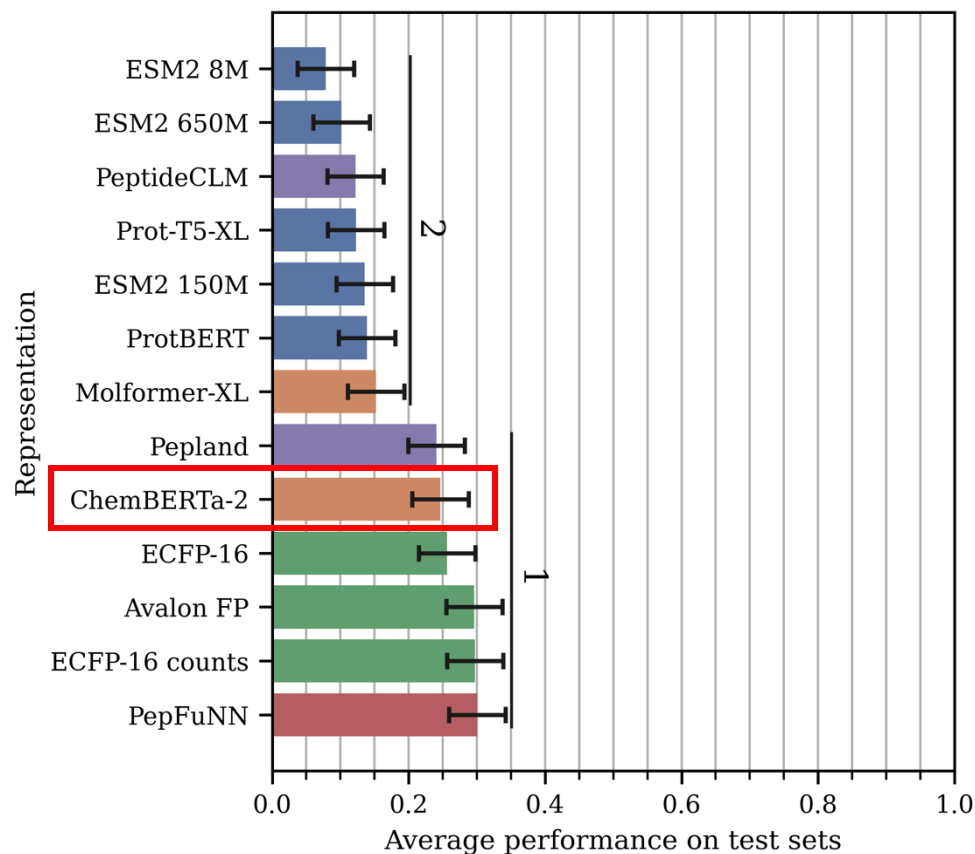
Synthetic peptides and peptidomimetics





Natural to synthetic extrapolation

- Average on 4 datasets.
- Statistical analysis is ANOVA with *post-hoc* Tukey's HSD



Conclusions

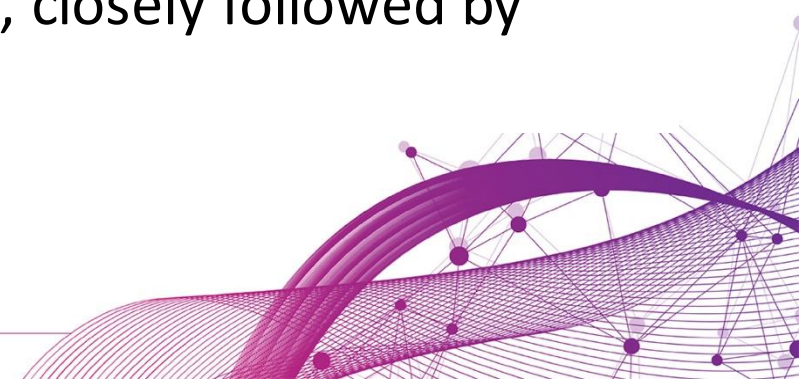
1. Natural to synthetic extrapolation is possible, but models are less reliable
2. ChemBERTa-2 appears to be the most versatile tool to work with peptides
3. Chemical and Peptide Fingerprints are robust options as well



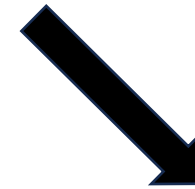


Conclusions

1. Dataset building (negative definition) and partitioning (train/test split) are key for proper model evaluation
2. Chemical fingerprints are better for partitioning natural and synthetic datasets than sequence alignment, contrary to standard practice
3. Natural to synthetic extrapolation is possible, but there is room for improvement. We release the benchmark datasets to make it easier for the community to improve
4. ChemBERTa-2 appears to be the most versatile tool, closely followed by chemical fingerprints



Contact info, papers, and
slides of the presentation



More information
and contact info



Partitioning, representation, and automation in canonical and non-canonical peptide modelling

Speaker: **Raúl Fernández-Díaz (PhD Candidate UCD – IBM Research)**

UCD: R. Cossio-Pérez, C. Agoni, D.C. Shields

IBM Research: T.L. Hoang, V. Lopez

Novo Nordisk: R. Ochoa

More information
and contact info

