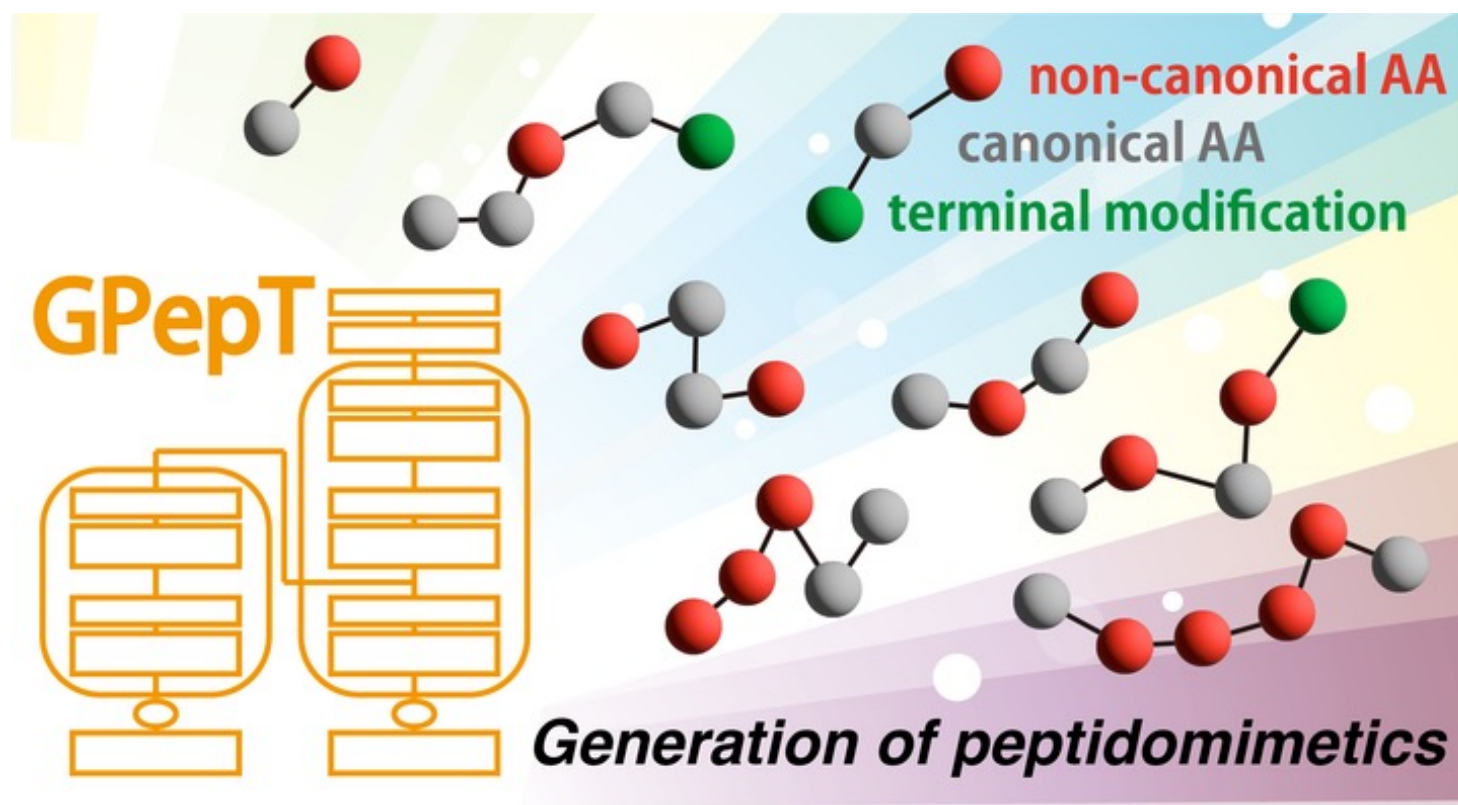
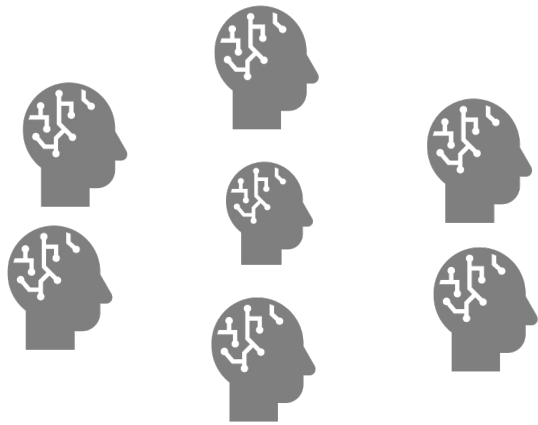


GPepT: A Foundation Language Model for Peptidomimetics Incorporating Noncanonical Amino Acids



Yuna Oikawa (D1), Tsuda Lab, UTokyo

Why **LLM** (and not other AI models) for bio ?

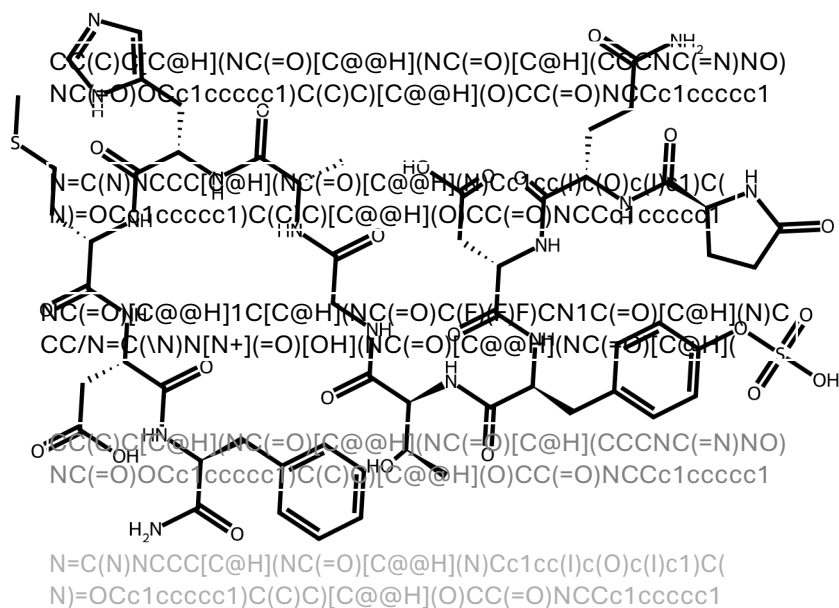


- **Easy** to use.
- Capability to **connect** with any data.

What should be done now?

- Feed LLMs with the “right” data.

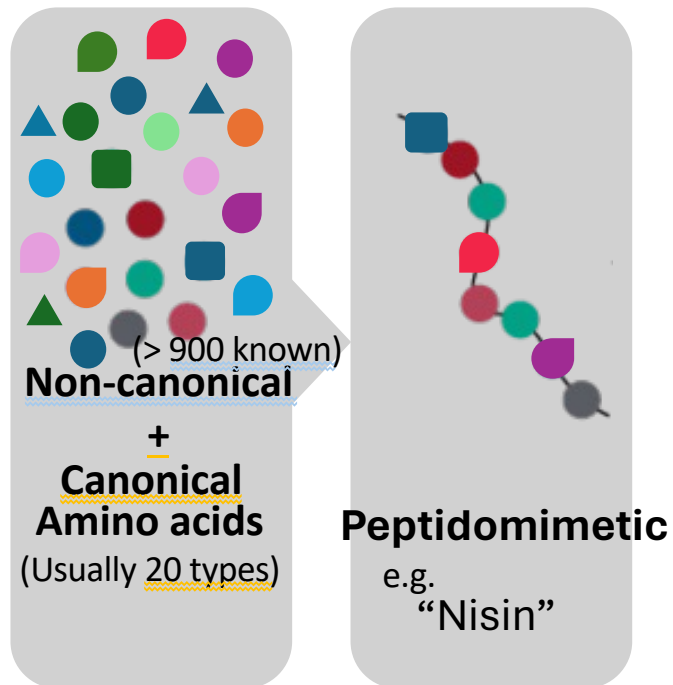
Consistent data:
Amino acid sequence



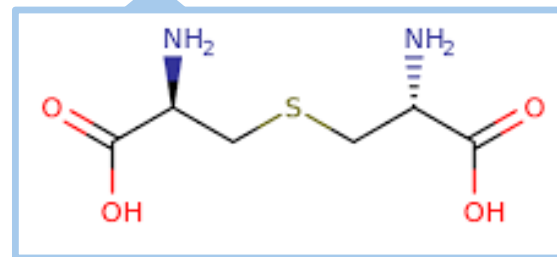
PQDPTGAHMDP

?QD?TGAHMD?

Non-canonical amino acids

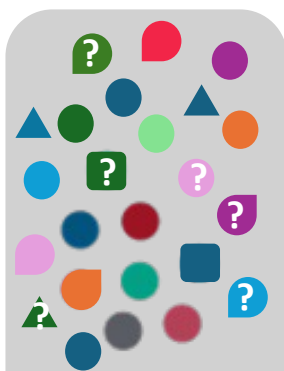


- e.g. Nisin :
 - Natural antimicrobial peptide
 - Uses lanthionine etc → unique structure



- Very few resistance reported

Peptidomimetics



**Non-canonical
Amino acids**
(> 900 types)

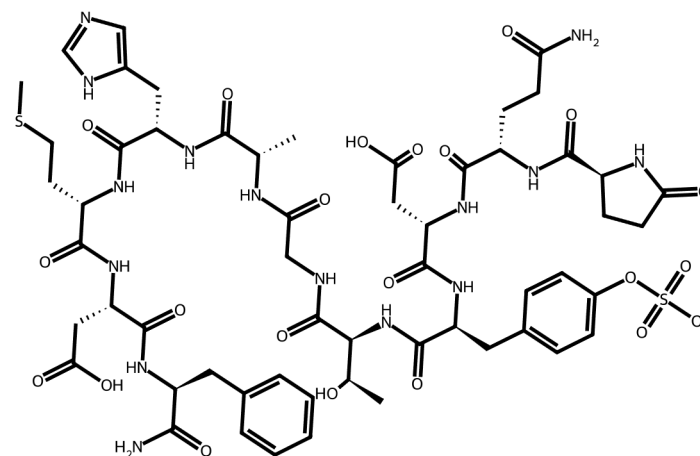


**Peptidomimetic
sequence**

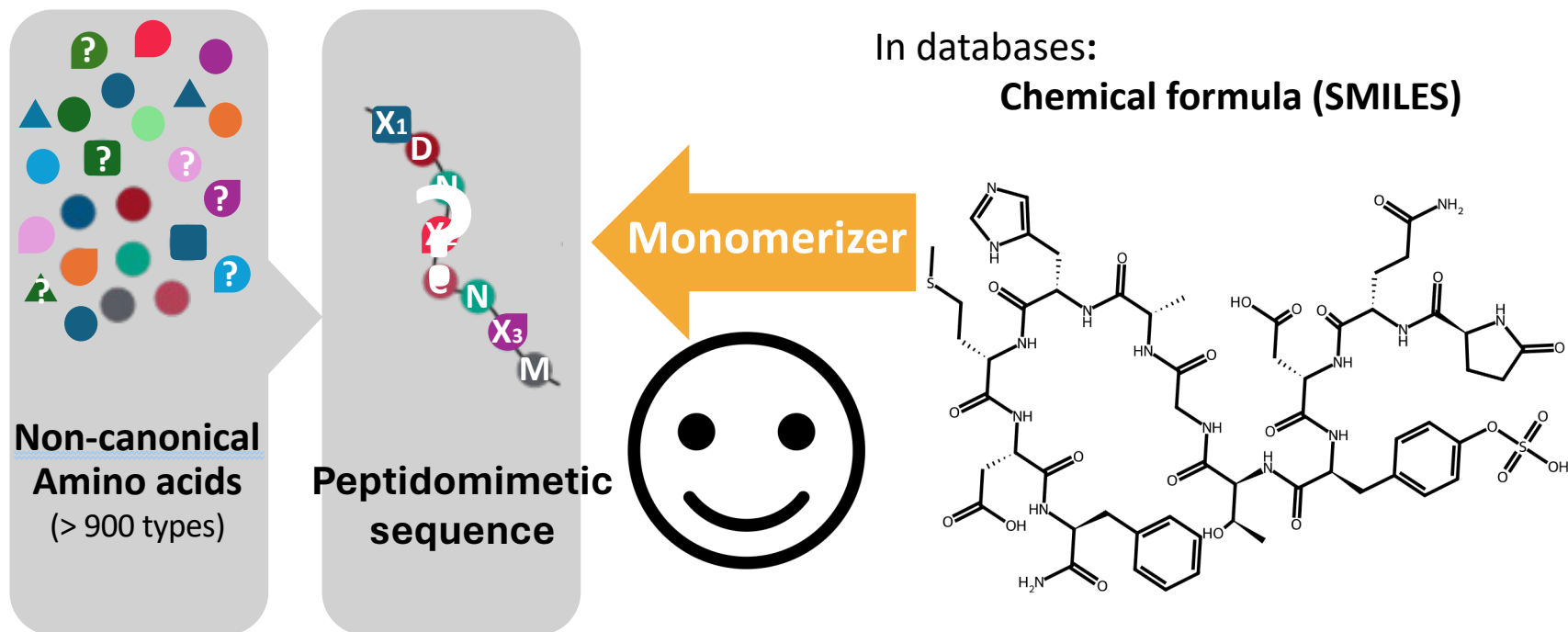


In databases:

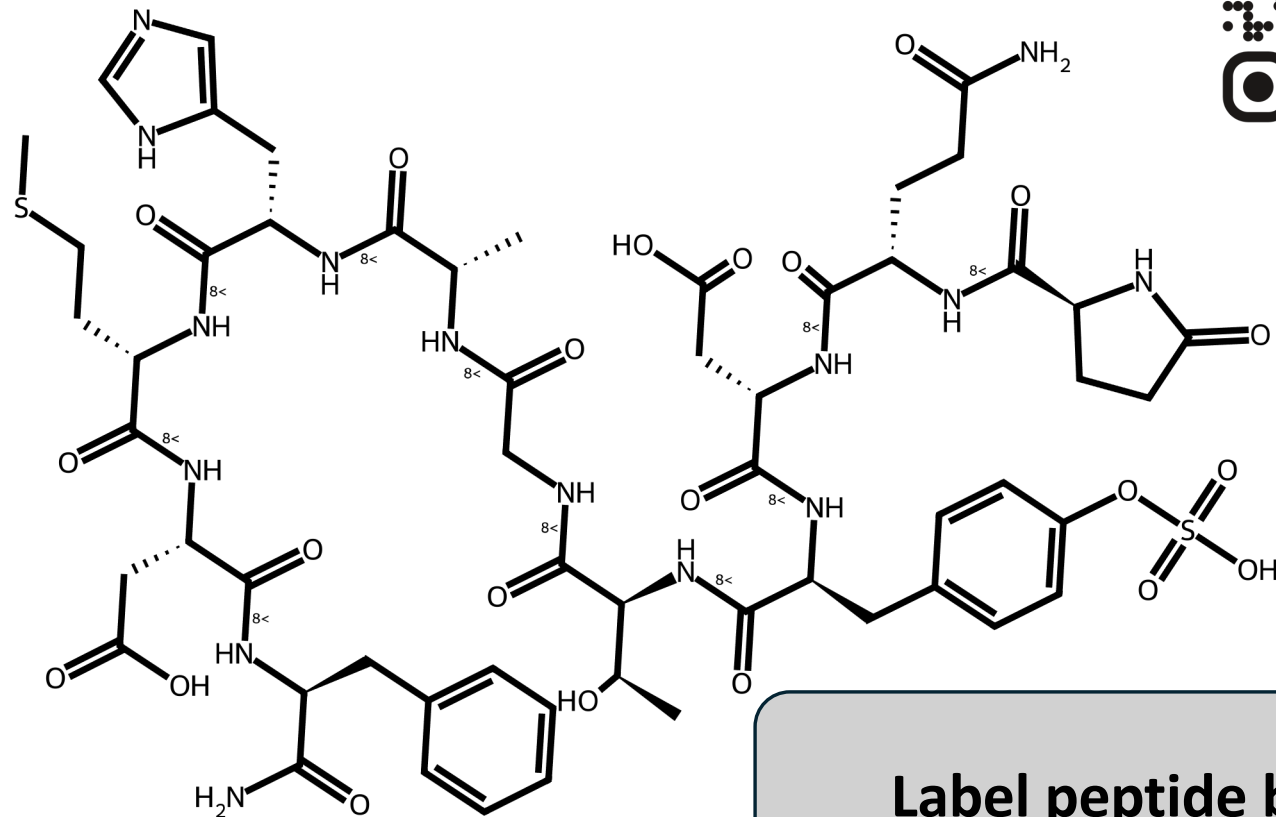
Chemical formula (SMILES)



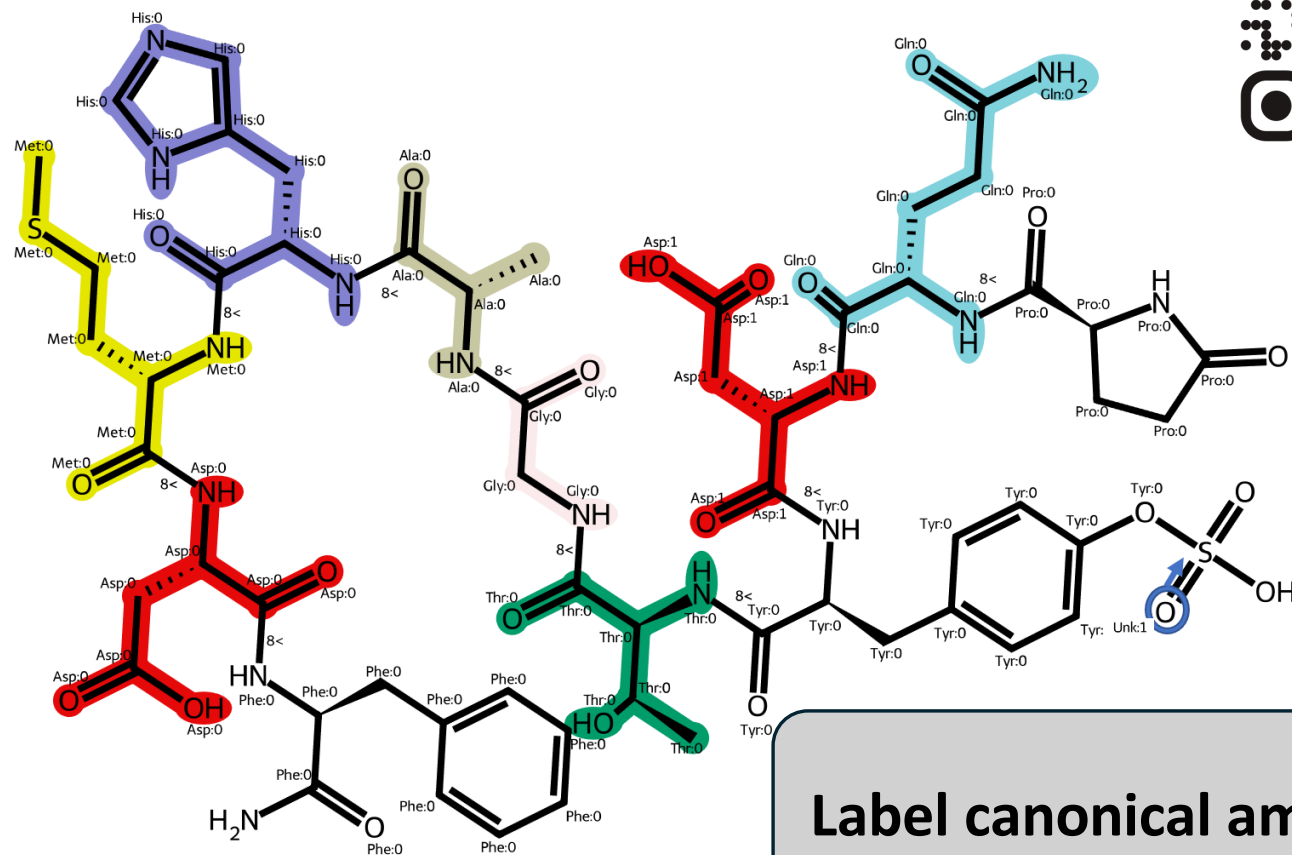
Peptidomimetics



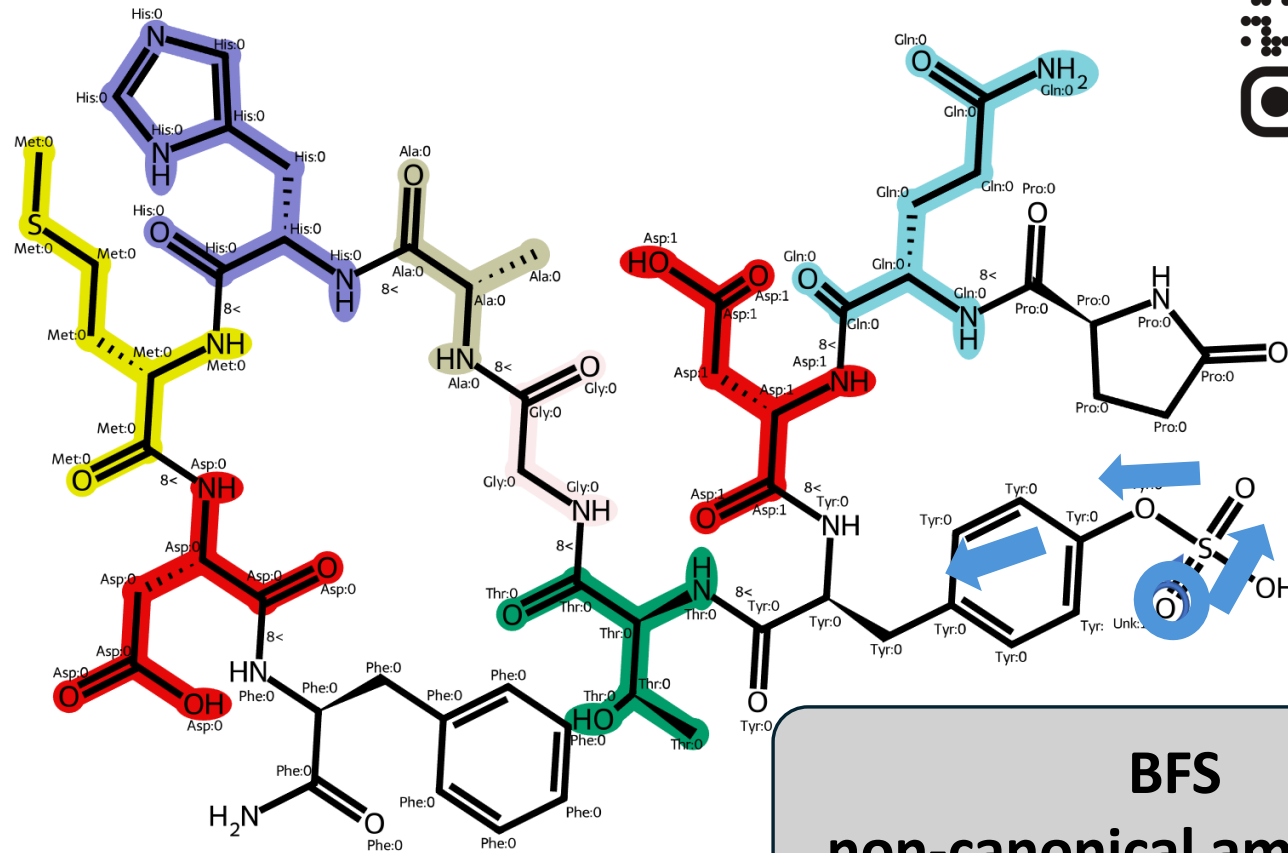
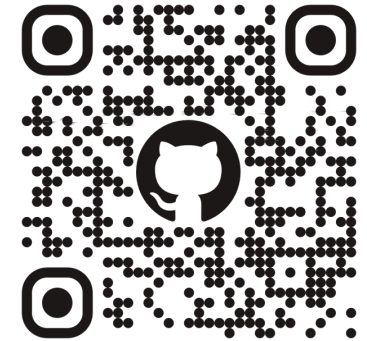
Monomerizer



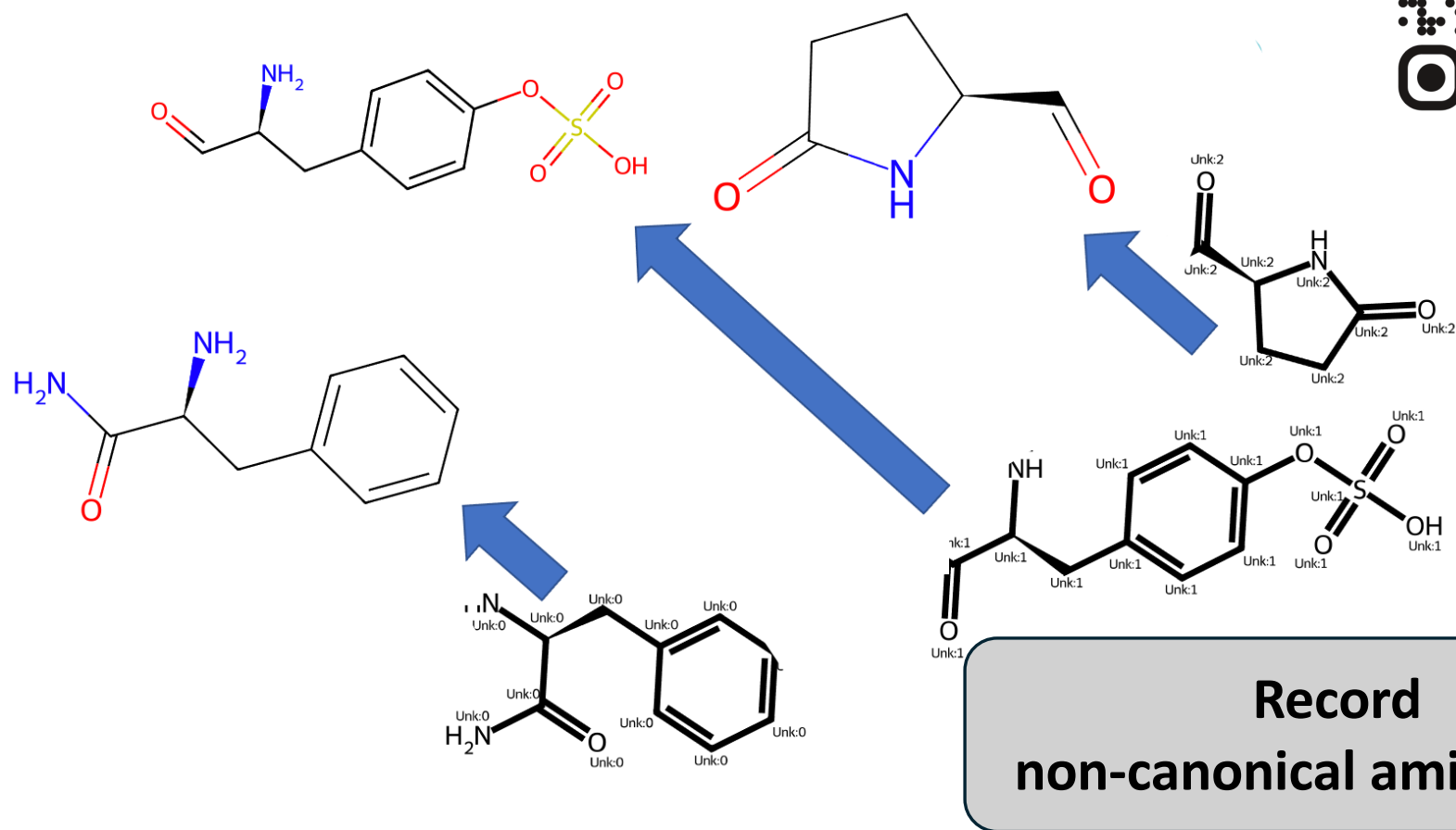
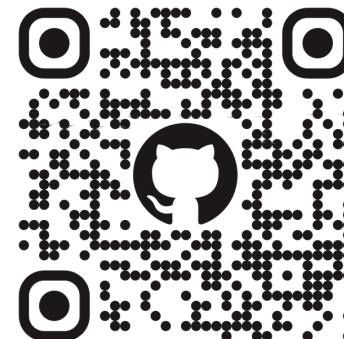
Label peptide bonds



Monomerizer

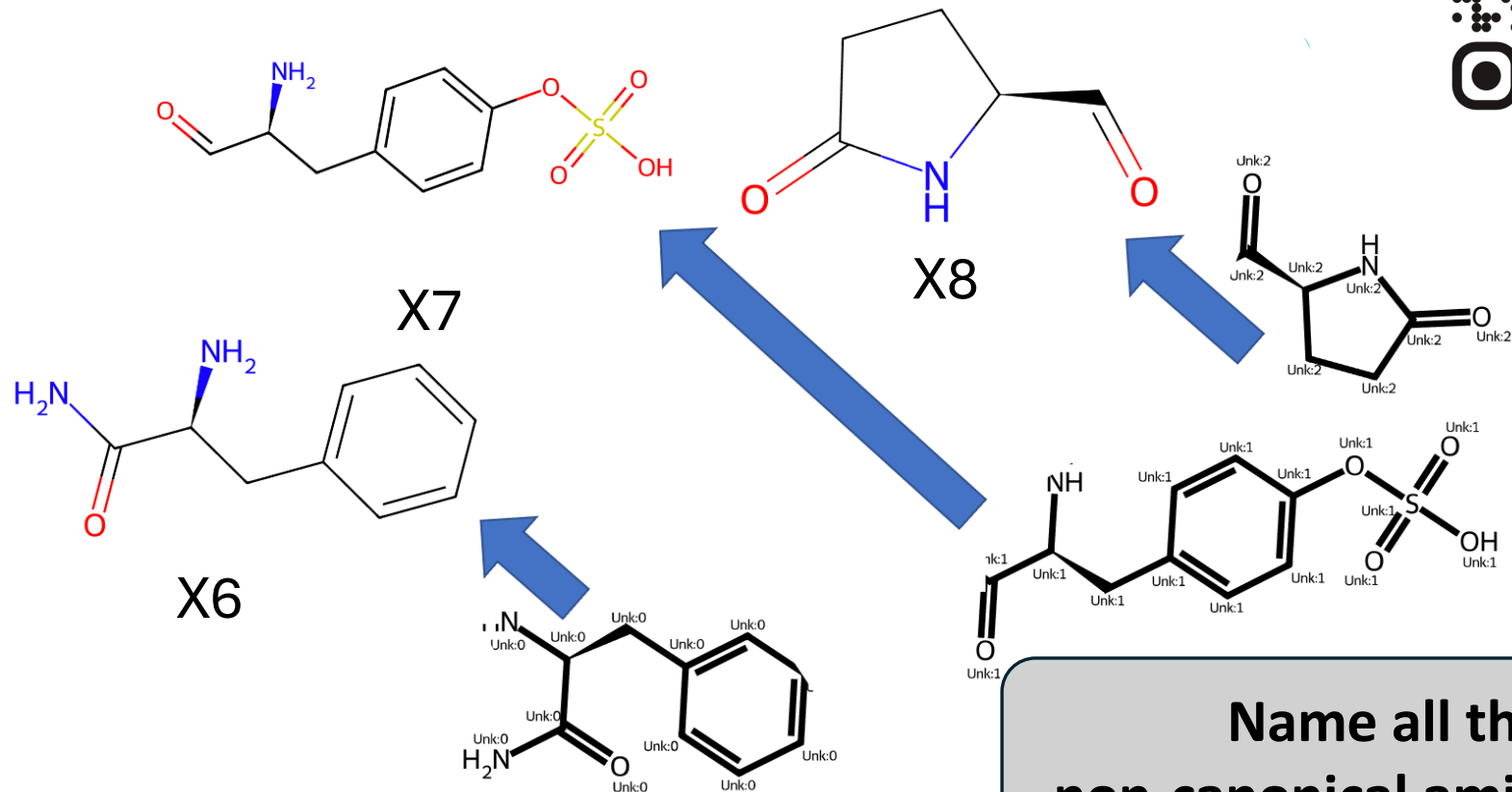
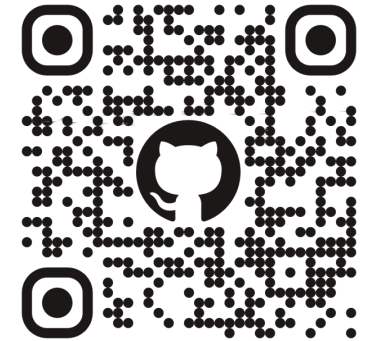


Monomerizer



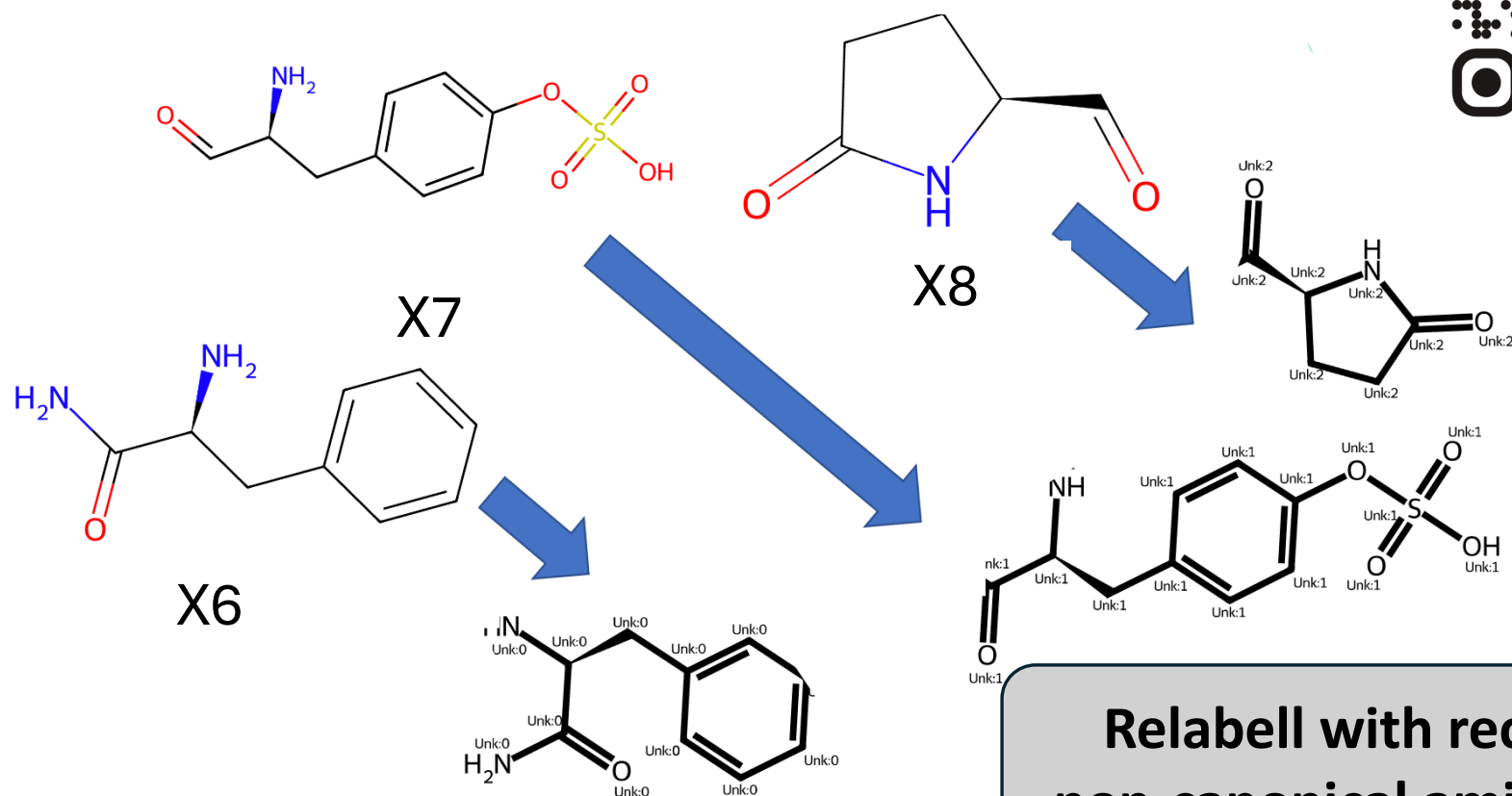
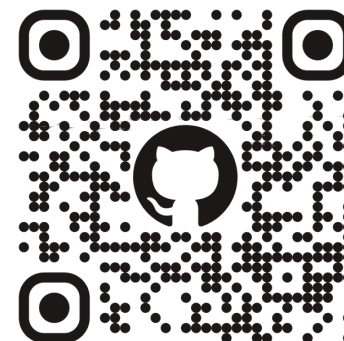


Monomerizer



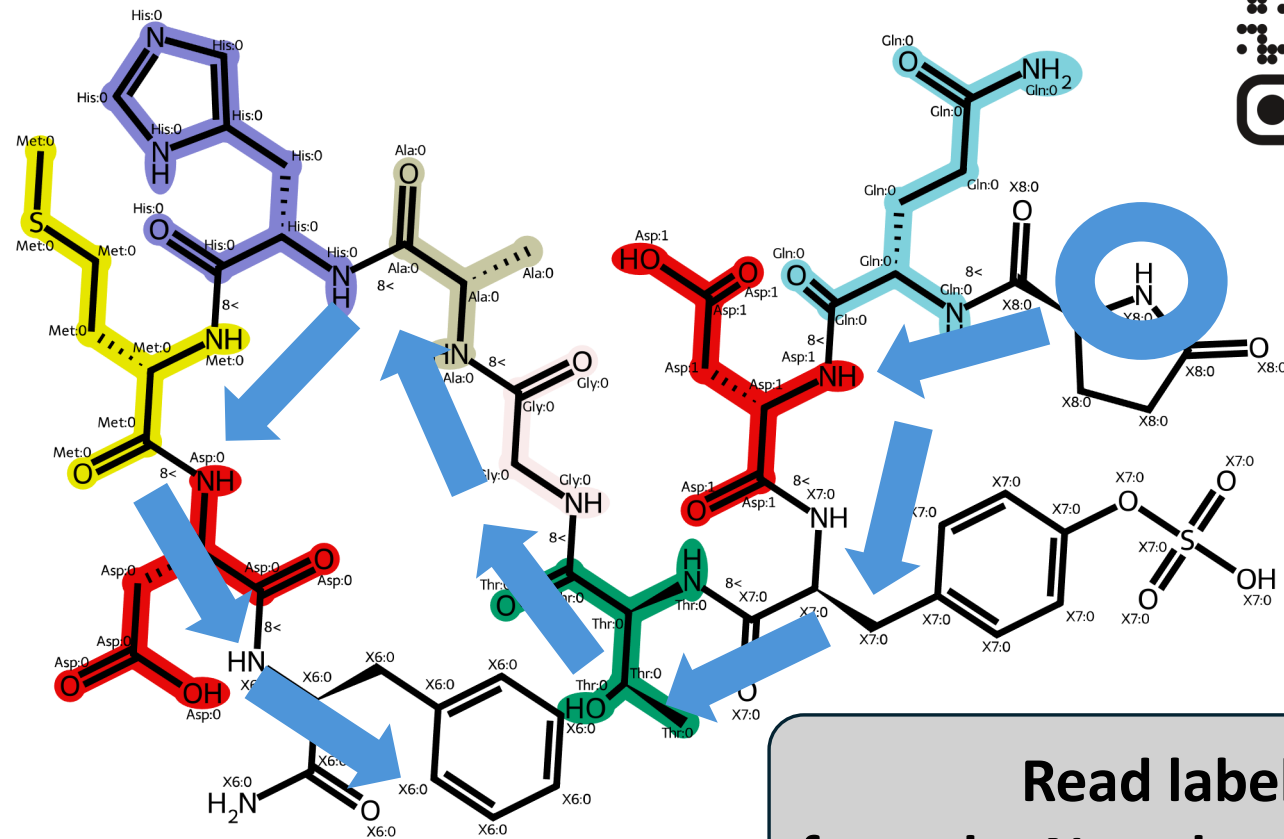
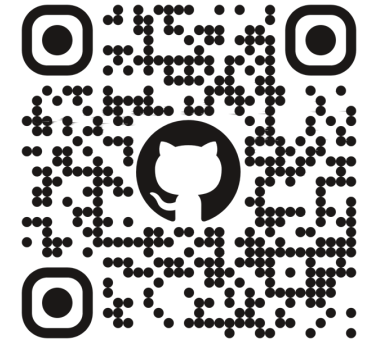
Name all the non-canonical amino acids

Monomerizer



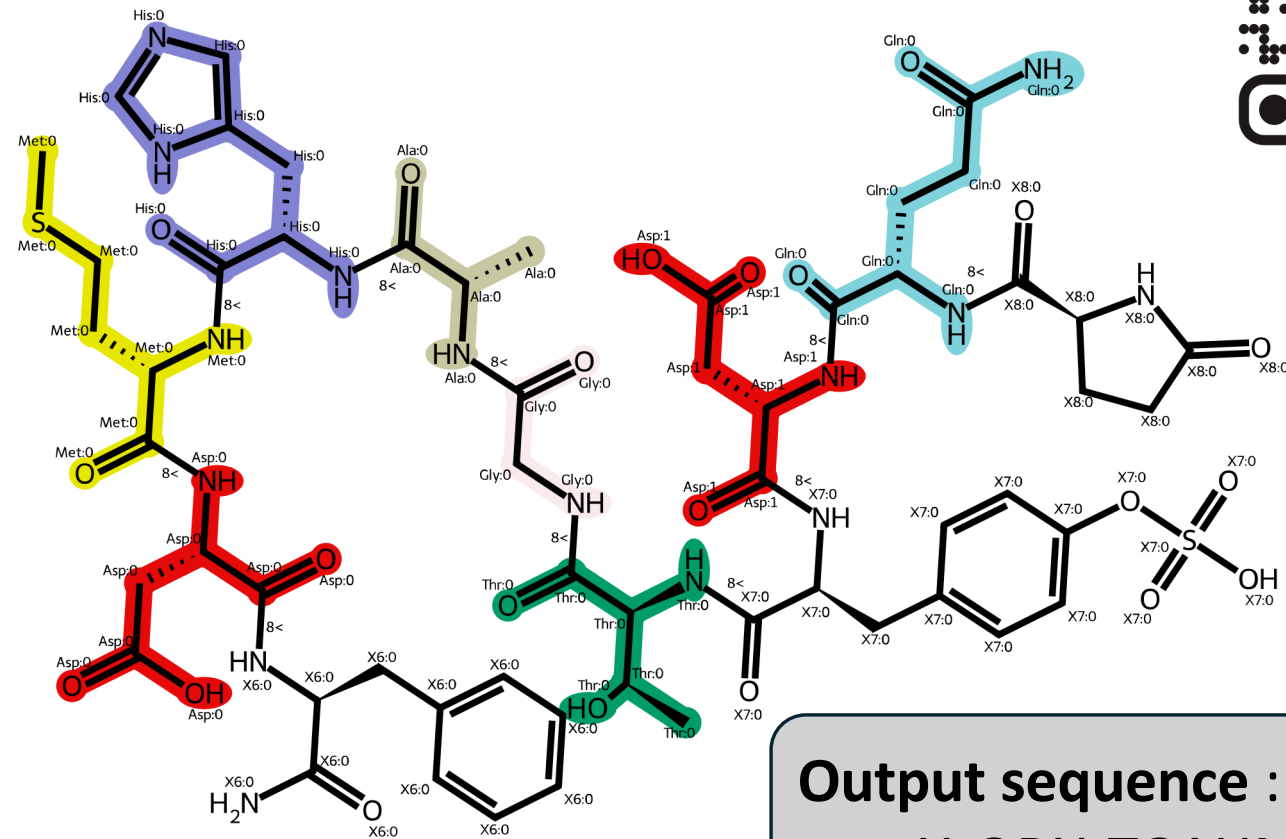
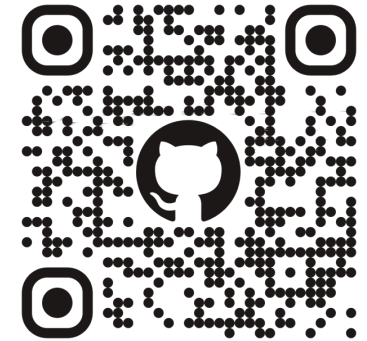
**Relabel with recorded
non-canonical amino acids**

Monomerizer



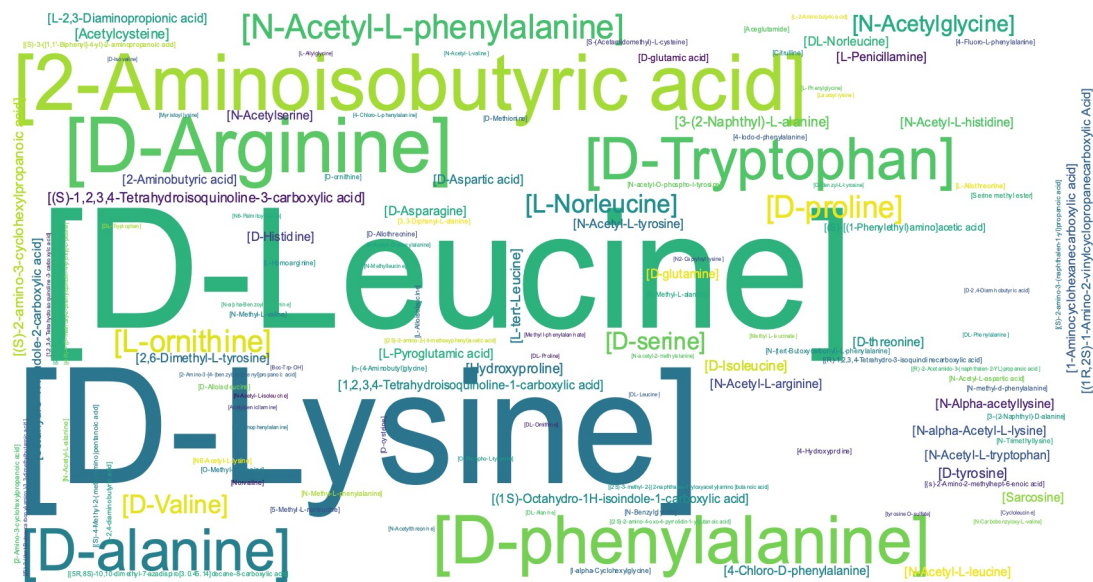
**Read labels
from the N to the C terminal**

Monomerizer

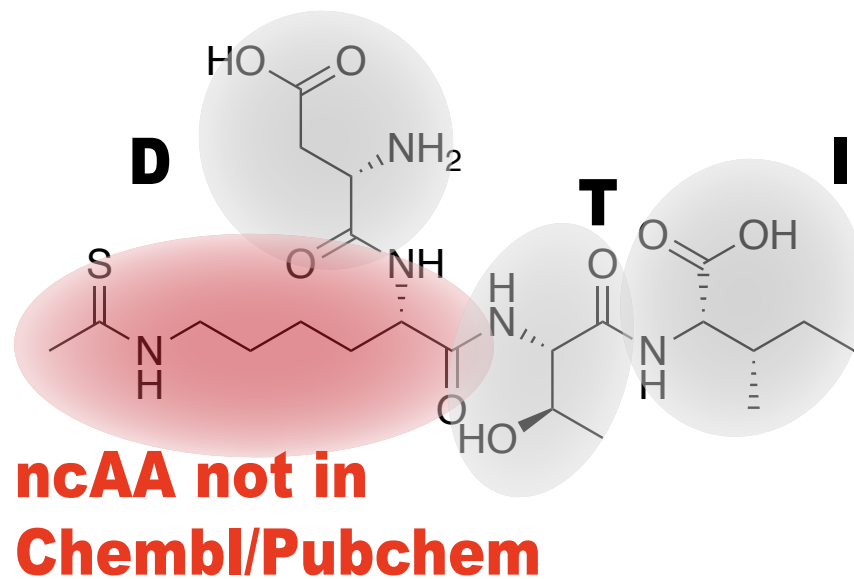


Output sequence :
X₈QDX₇TGAHMDX₆

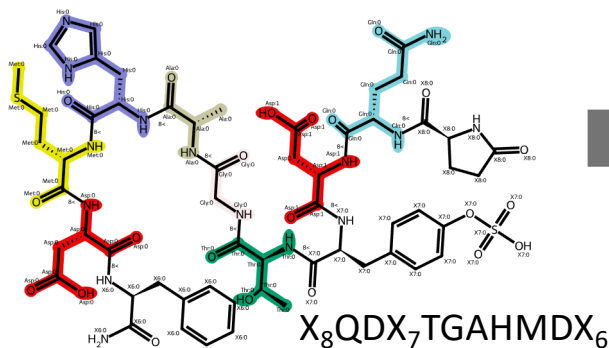
> 17,000 non-canonical amino acids



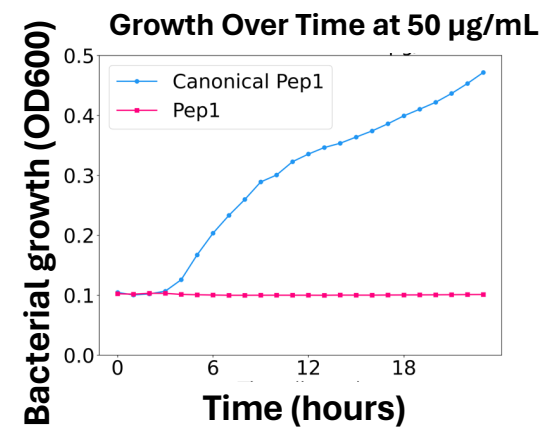
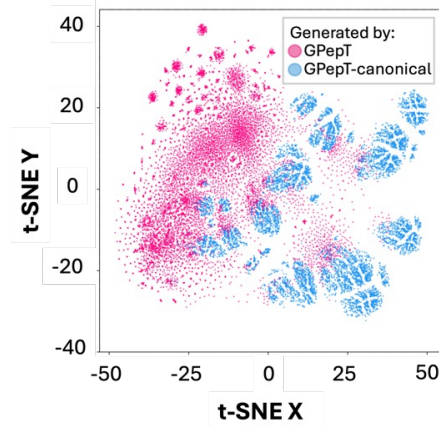
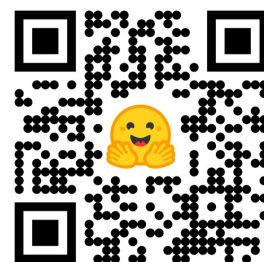
CHEMBL493160



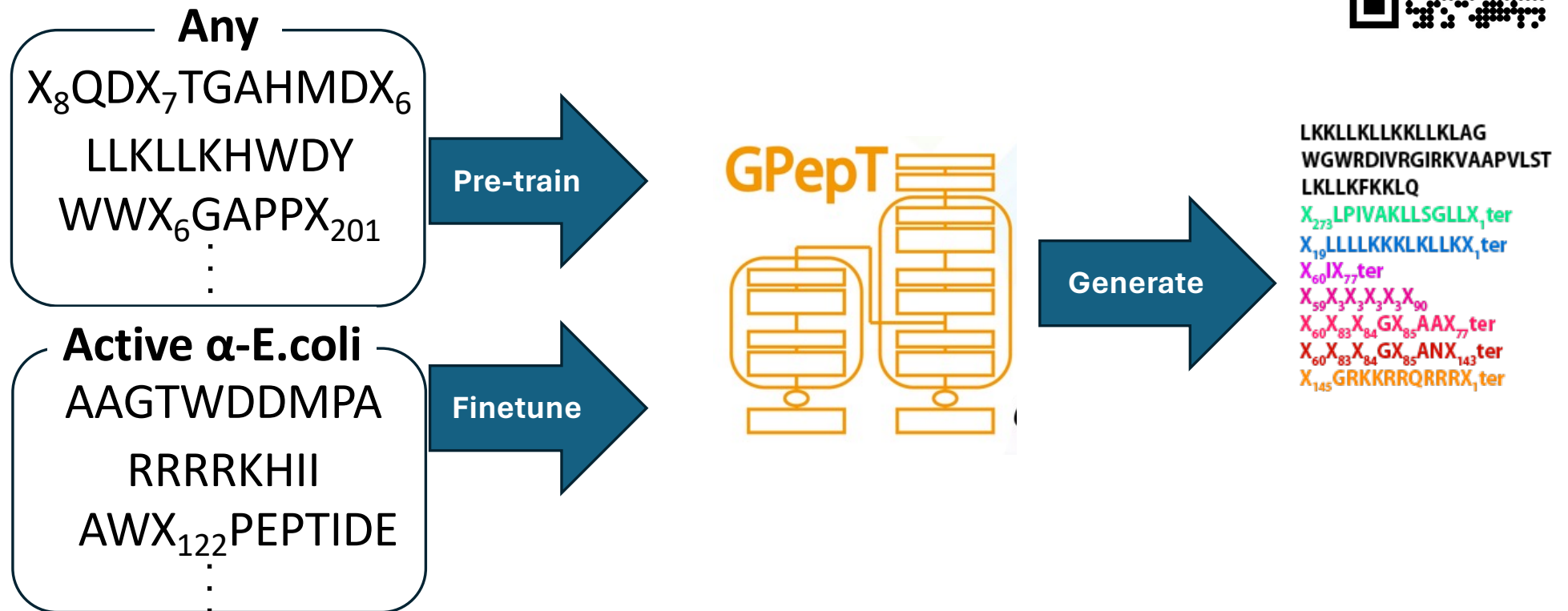
Overview



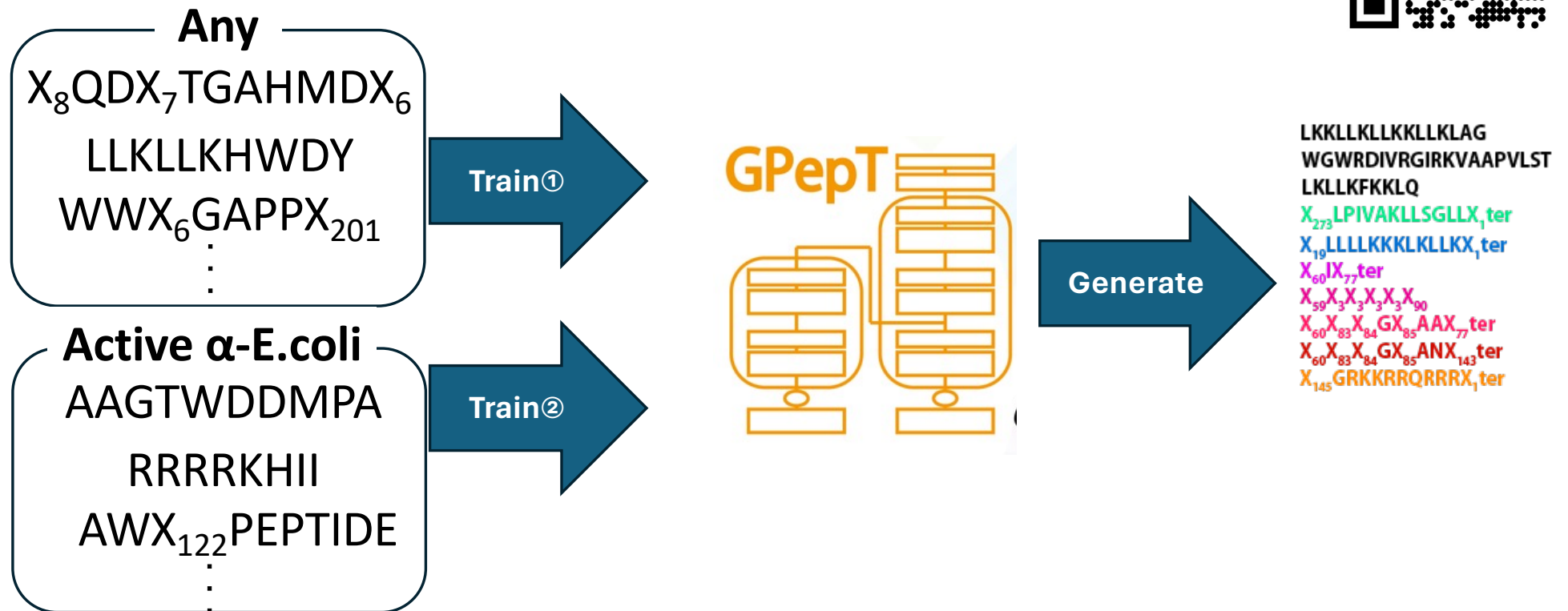
- 38,138 peptidomimetics
- 4,605 peptides



How to Make a Language Model Speak the Language of Peptides



How to Make a Language Model Speak the Language of Peptides



How to Make a Language Model Speak the Language of Peptides



1. Load the tokenizer, model, and config.

```
# Define the model path (e.g., "username/model-name")
model_repo = "openai-community/gpt2"

# Load the tokenizer, model, and config
tokenizer = AutoTokenizer.from_pretrained(model_repo)
model = AutoModelForCausalLM.from_pretrained(model_repo)
config = AutoConfig.from_pretrained(model_repo)
```

How to Make a Language Model Speak the Language of Peptides



2. Initialize model weights

```
model = AutoModelForCausalLM.from_config(config)
```

How to Make a Language Model Speak the Language of Peptides



3. Modify tokenizer

Playingyoyo/GPepT like 0

Text Generation PyTorch Safetensors gpt2 License: apache-2.0

Model card Files and versions Community Settings

main GPepT 1 contributor History: 67 commits + Contribute

File	Size	Upload	Time
Playingyoyo Update README.md c5fdb97 VERIFIED			about 2 months ago
.gitattributes	1.56 kB	Upload TOC.png	2 months ago
README.md	5.16 kB	Update README.md	about 2 months ago
TOC.png	582 kB	Upload TOC.png	2 months ago
config.json	918 Bytes	Upload model	7 months ago
generation_config.json	116 Bytes	Upload model	7 months ago
model.safetensors	3.1 GB	Upload model	8 months ago
pytorch_model.bin	3.13 GB	Upload model	8 months ago
requirements.txt	3.14 kB	Upload requirements.txt for run_clm.py	2 months ago
special_tokens_map.json	149 Bytes	Upload tokenizer	7 months ago
tokenizer.json	378 kB	Upload tokenizer	8 months ago
tokenizer_config.json	463 Bytes	Upload tokenizer	7 months ago

```
{
  "version": "1.0",
  "truncation": null,
  "padding": null,
  "added_tokens": [
    {
      "id": 0,
      "content": "<|endoftext|>",
      "single_word": false,
      "lstrip": false,
      "rstrip": false,
      "normalized": false,
      "special": true
    }
  ],
  "normalizer": null,
  "pre_tokenizer": {
    "type": "Split",
    "pattern": {
      "Regex": "(?=[A-Z])|(?=<\\|endoftext\\|>)"
    },
    "behavior": "Isolated",
    "invert": false
  },
  "post_processor": null,
  "decoder": null,
  "model": {
    "type": "WordLevel",
    "vocab": {
      "<|endoftext|>": 0,
      "[UNK]": 1,
      "X10000": 2,
      "X10001": 3,
      "X10002": 4,
      "X10003": 5,

```

How to Make a Language Model Speak the Language of Peptides



4. Adjust vocab size

```
# Load the tokenizer
tokenizer = AutoTokenizer.from_pretrained("Playingyoyo/GPePT")

# Resize the model's token embeddings to accommodate the new tokens
model.resize_token_embeddings(len(tokenizer))
```

How to Make a Language Model Speak the Language of Peptides



5. Pre-train

```
python run_clm.py --model_name_or_path Playingyoyo/GPepT \
--train_file path_to_train90.txt \
--validation_file path_to_val10.txt \
--tokenizer_name Playingyoyo/GPepT \
--do_train \
--do_eval \
--output_dir ./output \
--learning_rate 1e-5
```

No GPU
required

How to Make a Language Model Speak the Language of Peptides



6. Finetune

```
python run_clm.py --model_name_or_path Playingyoyo/GPepT \
--train_file path_to_train90.txt \
--validation_file path_to_val10.txt \
--tokenizer_name Playingyoyo/GPepT \
--do_train \
--do_eval \
--output_dir ./output \
--learning_rate 1e-5
```

No GPU
required

How to Make a Language Model Speak the Language of Peptides



7. Generate

```
# Generate sequences (expressed in tokens, average ~4 amino acids per token)
sequences = GPepT("<|endoftext|>",
                  max_length=25,
                  do_sample=True,
                  top_k=950,
                  repetition_penalty=1.5,
                  num_return_sequences=5,
                  eos_token_id=0)

# Print generated sequences
for seq in sequences:
    print(seq['generated_text'])
```

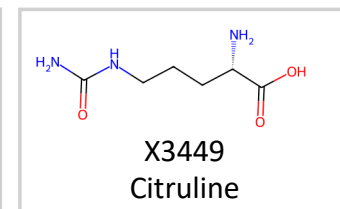
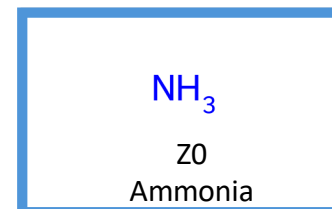
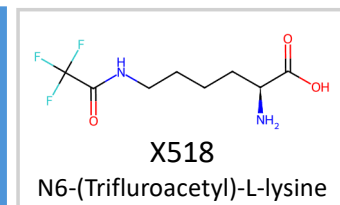
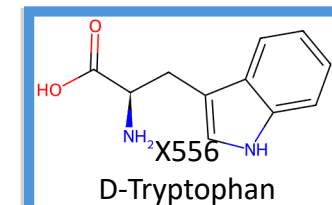
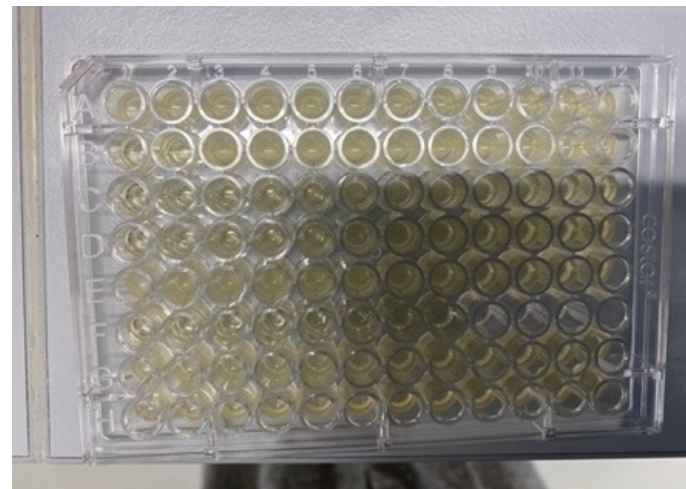
No GPU
required

Experimental validation

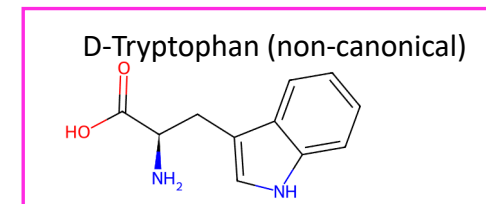
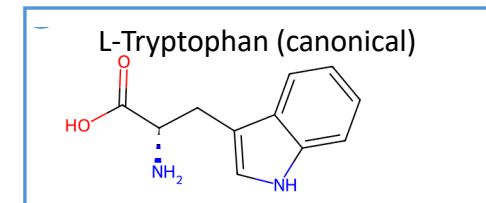
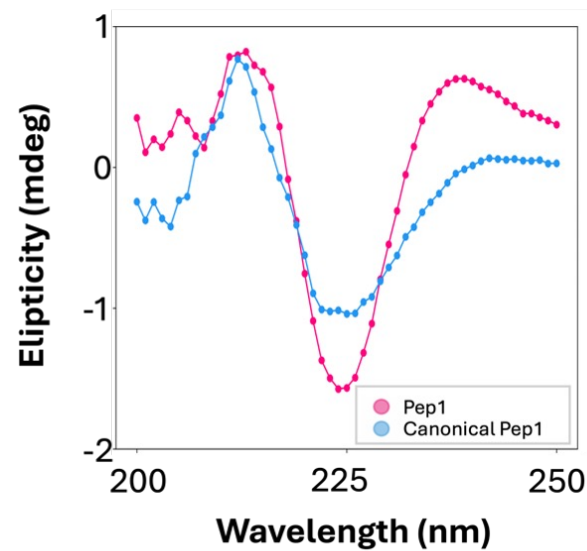
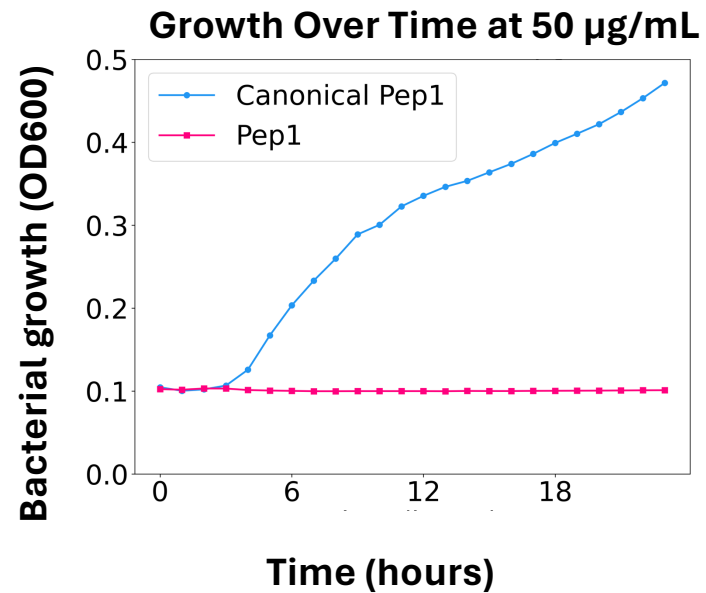


MIC=50 µg/mL

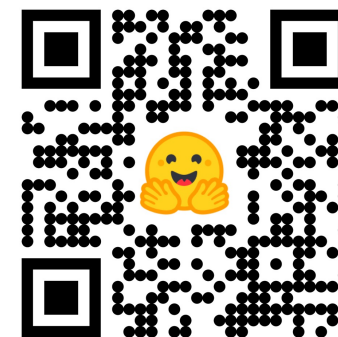
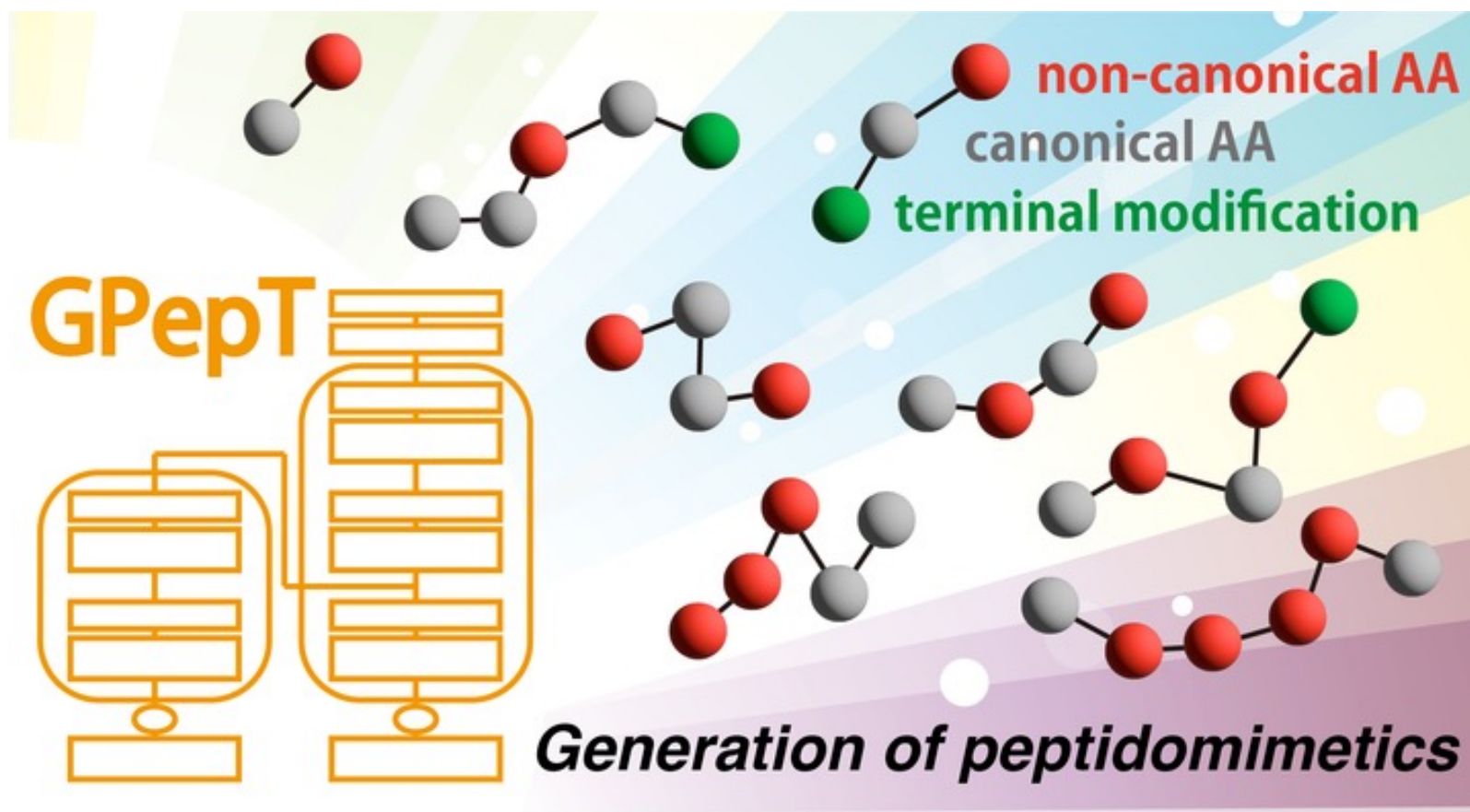
Name	Sequence	Monomers
Pep1	X556WX556WKZ0	X556=D-Tryptophan, Z0=Ammonia
Pep2	X556WWWWWWWWWWWWWW	X556=D-Tryptophan
Pep3	LQKYRRVRGGRX518F	X518=N6-(Trifluoroacetyl)-L-lysine
Pep4	QGRKQGRRX518	X518=N6-(Trifluoroacetyl)-L-lysine
Pep5	X3449WWWWWR	X3449=Citrulline



Experimental validation



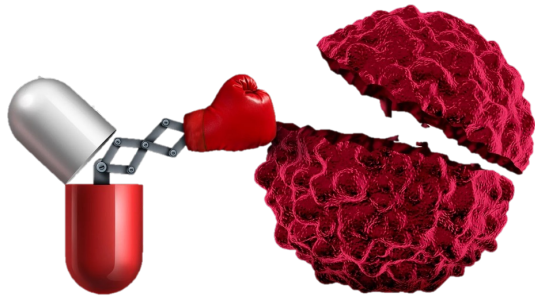
First Peptidomimetic Language Model



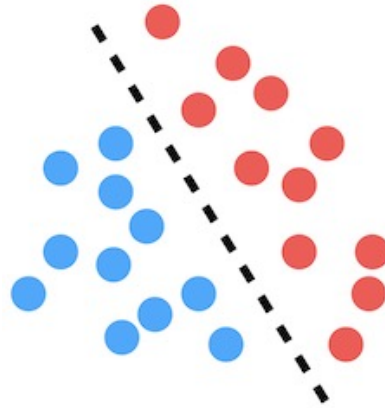
Future Prospects①



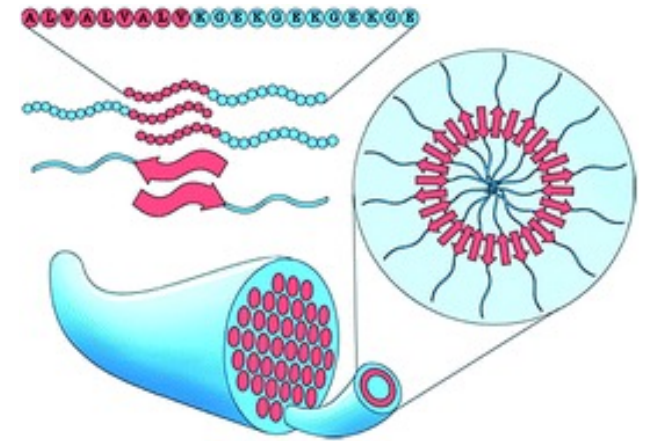
Anti-cancer
peptidomimetics generation



Toxicity classifier



Fiber development



Future Prospects②



ChatGPT ▾

↑ Share ...

when

X7681 = Morpholine,

Z81 = Ethyl carbazate,

Z0 = Amide,

X1132 = 3,5-Diiodo-L-tyrosine,

X415 = 3,4-Diethoxyaniline and

X2326 = Glycine, N-(3-iodobenzoyl)-,

which of the peptidomimetics sequences below are most likely to be active against *E. coli* ?

X7681VZ81

X1132RZ0

X415X2326Z0

1. X7681VZ81

- **X7681 = Morpholine**: A small, polar, heterocyclic compound. Not inherently antimicrobial, but found in drug scaffolds.
- **V = Valine**: Hydrophobic amino acid; may aid membrane interaction.
- **Z81 = Ethyl carbazate**: Can form hydrogen bonds; may enhance stability and interaction.
→ Moderate polarity with potential hydrogen bonding and slight hydrophobicity. May cross membranes, but lacks strong antimicrobial motifs.

⚖️ **Verdict**: Possibly weak activity.

✓ Conclusion:

Most likely to be active against *E. coli*:

→ **X1132RZ0** (Sequence #2)

This contains a **halogenated aromatic + Arginine**, a classic combination for antimicrobial activity via membrane disruption.

First Peptidomimetic Language Model

