

RESEARCH ARTICLE

[View Article Online](#)
[View Journal](#) | [View Issue](#)Cite this: *RSC Med. Chem.*, 2024, 15, 2030Received 8th March 2024,
Accepted 19th April 2024

DOI: 10.1039/d4md00159a

rsc.li/medchemCan large language models predict antimicrobial peptide activity and toxicity?[†]Markus Orsi  and Jean-Louis Reymond *

Antimicrobial peptides (AMPs) are naturally occurring or designed peptides up to a few tens of amino acids which may help address the antimicrobial resistance crisis. However, their clinical development is limited by toxicity to human cells, a parameter which is very difficult to control. Given the similarity between peptide sequences and words, large language models (LLMs) might be able to predict AMP activity and toxicity. To test this hypothesis, we fine-tuned LLMs using data from the Database of Antimicrobial Activity and Structure of Peptides (DBAASP). GPT-3 performed well but not reproducibly for activity prediction and hemolysis, taken as a proxy for toxicity. The later GPT-3.5 performed more poorly and was surpassed by recurrent neural networks (RNN) trained on sequence-activity data or support vector machines (SVM) trained on MAP4C molecular fingerprint-activity data. These simpler models are therefore recommended, although the rapid evolution of LLMs warrants future re-evaluation of their prediction abilities.

Introduction

Antimicrobial peptides (AMPs) have gained significant attention in the field of drug discovery due to their potential therapeutic applications in the fight against antimicrobial resistance.^{1–3} However, the vast number of possible peptide sequences and their complex structure–activity relationship landscape mean that it is difficult to rationally design peptides with the desired biological activity, in particular tuning their activity *versus* toxicity to human cells, which is often measured as hemolysis of human red blood cells.^{4,5}

To address this issue, several machine-learning models have been developed for the *de novo* design of antimicrobial peptides.^{6–21} Because property prediction from a peptide sequence can be framed as a natural language processing problem, many of these models use architectures specifically designed for language processing tasks.^{22–24} Furthermore, the emergence of large language models (LLMs), such as OpenAI's GPT models,²⁵ has opened new possibilities for leveraging powerful language processing capabilities in drug discovery applications. Recent attempts by Jablonka *et al.* to explore the capabilities of GPT-3 for predicting properties of small molecules in various applications have shown that GPT-3 was able to perform comparably or even outperform conventional statistical models, particularly in the low data regime.²⁶ There also have been successful efforts into

augmenting LLM capabilities to tackle tasks related to small molecule chemistry in the areas of organic synthesis, drug discovery, and materials design.^{27–30} Hereby, the models mainly orchestrate a set of tools to solve chemistry tasks starting from a natural language prompt.^{31–33} However, to the best of our knowledge LLMs have not been implemented to predict the bioactivity of peptides yet.

In this study, we aimed to compare GPT models fine-tuned on antimicrobial peptide sequence data with models that have been previously used to predict antimicrobial activity and hemolysis of peptide sequences.^{13,14} Alongside evaluating the performance of the fine-tuned GPT models, we also seek to explore the advantages and disadvantages they offer in terms of time and cost effectiveness. Furthermore, we compare the performance of models trained on amino acid sequences to a support-vector machine (SVM) trained on the MAP4C fingerprint.³⁴

Methods

Datasets

The datasets used in this study were peptide sequences with annotated antimicrobial and hemolytic activity collected from the Database of Antimicrobial Activity and Structure of Peptides (DBAASP).^{13,35} Sequences exhibiting an activity measure below 10 mM, equivalent to 10 000 nM or 32 mg mL^{−1}, against at least one of the selected target organisms *P. aeruginosa*, *A. baumannii*, or *S. aureus* were categorized as active. Conversely, sequences with activity measures exceeding 10 mM, 10 000 nM, or 32 mg mL^{−1} against all of these targets were categorized as inactive. When available,

Department of Chemistry, Biochemistry and Pharmaceutical Sciences, University of Bern, Freiestrasse 3, 3012 Bern, Switzerland. E-mail: jean-louis.reymond@unibe.ch

[†] Electronic supplementary information (ESI) available. See DOI: <https://doi.org/10.1039/d4md00159a>



activity against human erythrocytes was utilized to classify sequences as either hemolytic or non-hemolytic. Concentrations were standardized to mM, and sequences causing less than 20% hemolysis at concentrations equal to or above 50 mM were categorized as non-hemolytic and flagged accordingly. Sequences inducing more than 20% hemolysis were classified as hemolytic, irrespective of concentration. The dataset used for the classification tasks contained 9548 (7160 training/2388 validation) sequences with annotated antimicrobial activity, of which 2262 (1723 training/539 validation) sequences had additional hemolytic activity annotations. To test models in low data regimes, we randomly selected subsets from the original training sets, representing approximately 20% and 2% of the original activity set, and approximately 10% of the original hemolysis set. All datasets are further described in Table 1. To ensure consistency, we maintained the same training and test split for all initial evaluations. For the detailed study, we used the same 5-fold cross-validation sets.

Models

As reference models, we used our previously reported naïve Bayes (NB), support vector machine (SVM), random forest (RF), and recurrent neural network (RNN) classifiers trained on the same data.¹³ We furthermore trained two additional SVM models on alternative representations of peptide sequences: one utilizing the MAP4C fingerprint³⁴ with a custom Jaccard kernel, and another using predicted fraction of helical residues and hydrophobic moment with a linear kernel. Fraction of helical residues were predicted using SPIDER3.³⁶ Hydrophobic moment was computed using the method of Eisenberg *et al.*³⁷

To explore the potential of GPT-3 models for antimicrobial and hemolytic activity classification, we performed fine-tuning of the Ada, Babbage, and Curie models which were accessible through the OpenAI API (v0.28.0, accessed between 25.05.2023 and 01.06.2023). The fine-tuning process involved training each model using the full, 20% and 2% sets for activity classification and the full and 10% set for the hemolysis classification. In the later evaluation with the more advanced LLM GPT-3.5 Turbo, fine-tuning was also performed *via* OpenAI's Python API (v1.11.1), following the provided guidelines, but we restricted ourselves to the full

model. The utilized fine-tuning datasets contained a system role ("predicting antimicrobial activity/hemolysis from an amino acid sequence"), a user message (peptide sequence formatted as "SEQUENCE ->"), and a system message ("0" for negative labels and "1" for positive labels).

Metrics

All models were evaluated using five commonly accepted performance metrics: ROC AUC, accuracy, precision, recall and F1. Metrics were either calculated using the scikit-learn (v1.4.0) Python (v3.12.1) package (reference models and GPT-3.5) or directly obtained from the OpenAI platform after fine-tuning was completed (for all GPT-3 models).

ROC AUC (receiver operating characteristic area under the curve). The ROC AUC measures the area under the receiver operating characteristic curve, which plots the true positive rate (sensitivity) against the false positive rate. A higher ROC AUC value (ranging from 0 to 1) indicates better discrimination and predictive performance of the model.

Accuracy. Accuracy measures the overall correctness of the model's predictions, calculating the ratio of correctly classified instances to the total number of instances. It provides a general understanding of the model's performance but can be misleading in imbalanced datasets.

$$\text{Accuracy} = \frac{\text{TP} + \text{TN}}{\text{TP} + \text{FN} + \text{TN} + \text{FP}}$$

Precision. Precision measures the proportion of true positives out of all predicted positives. It focuses on the model's ability to avoid false positives.

$$\text{Precision} = \frac{\text{TP}}{\text{TP} + \text{FP}}$$

Recall. Recall measures the proportion of true positives out of all actual positives. It represents the model's ability to identify positive instances accurately.

$$\text{Recall} = \frac{\text{TP}}{\text{TP} + \text{FN}}$$

F1 score. F1 is the harmonic mean of precision and recall. It provides a balanced measure that considers both precision and recall.

$$\text{F1} = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}}$$

Table 1 Sizes and composition of the datasets used in the present study. Datasets are available at https://github.com/reymond-group/LLM_classifier

Name	Size	# positive class	# negative class
Activity training	7160	3580	3580
Activity training 20%	1400	701	699
Activity training 2%	140	74	66
Activity validation	2388	1194	1194
Hemolysis training	1723	717	1006
Hemolysis training 10%	170	65	105
Hemolysis validation	539	226	313

Results and discussion

Model screening

Starting from the DBAASP dataset of 9548 peptide sequences annotated with antibacterial activity and 2262 peptide sequences annotated with hemolysis effect, we had previously



Table 2 Performance metrics of all models tested on antimicrobial activity and hemolysis classification. The best value for each metric is highlighted in bold. NB: naïve Bayes, RF: random forest, SVM: support vector machine, RNN: recurrent neural network, MAP4C: chiral MinHashed atom-pair fingerprint of diameter 4, GPT: generative pre-trained transformer

Model	ROC AUC	Accuracy	Precision	Recall	F1
NB act.	0.55	0.55	0.59	0.32	0.42
RF act.	0.81	0.71	0.7	0.75	0.73
SVM act.	0.75	0.68	0.68	0.68	0.68
RNN act.	0.84	0.76	0.74	0.8	0.77
Features SVM act.	0.65	0.65	0.66	0.62	0.64
MAP4C SVM act.	0.8	0.8	0.79	0.83	0.8
GPT-3.5 Turbo act.	0.68	0.68	0.62	0.93	0.75
NB hem.	0.58	0.56	0.48	0.76	0.59
RF hem.	0.8	0.77	0.81	0.6	0.69
SVM hem.	0.69	0.73	0.72	0.58	0.65
RNN hem.	0.87	0.76	0.7	0.76	0.73
Features SVM hem.	0.62	0.63	0.57	0.5	0.54
MAP4C SVM hem.	0.83	0.83	0.76	0.85	0.8
GPT-3.5 Turbo hem.	0.65	0.69	0.72	0.43	0.54

evaluated NB, RF, SVM and RNN models, and found the latter to perform best for predicting both activity and hemolysis from sequence data.^{13,14} For additional reference, we trained an SVM on the fraction of helical residues and the hydrophobic moment, two properties commonly known to correlate with antimicrobial activity, as well as another SVM on MAP4C, a molecular fingerprint that can reliably encode large molecules such as natural products and peptides including their chirality,³⁴ a parameter which we considered important since our data listed sequences containing both L- and D-amino acids.

Aiming to test how LLMs perform in predicting antimicrobial activity and hemolysis, we first fine-tuned and evaluated GPT-3 Ada, Babbage, and Curie models. As discussed in our preprint, these models performed slightly better than the reference models, and even provided good performances when trained in low data regime (20% and 2% of full data). However, these models were later deprecated by OpenAI and their performance cannot be reproduced. We therefore discuss herein only the results obtained with the more recent GPT-3.5 model, in comparison with the reference models.

For both, prediction of antimicrobial activity and prediction of hemolysis, the top-performing models were the MAP4C SVM and the RNN model trained on sequence data, the latter being the best performer in our original work (Table 2).¹³ The performances for both models were in a similar range, although the RNN displayed a notably higher ROC-AUC in both tasks. GPT-3.5 displayed the highest recall performance among the activity models, indicative of the model's tendency to overly favor positive predictions, potentially leading to increased false positive predictions. On the other hand, the features SVM trained only on helicity and hydrophobic moment did not perform significantly above background, and was later used as a negative control model.

Model comparison

Following the initial model screening, we aimed to validate our findings through a more robust approach: a 5-fold cross-validation involving GPT-3.5, the MAP4C SVM, the RNN, and finally the features SVM as negative control. For this purpose, we generated five data splits and conducted predictions anew.

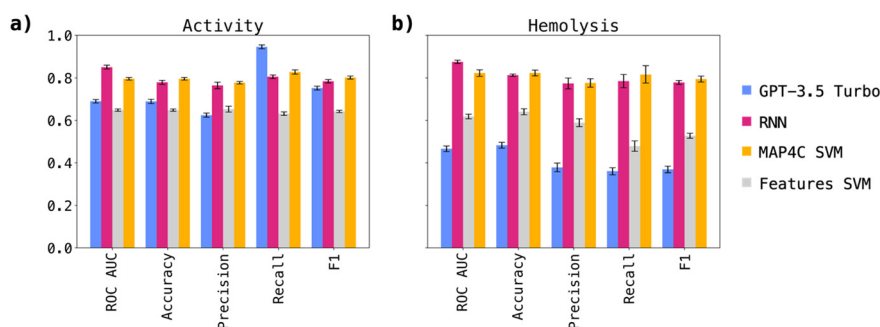


Fig. 1 Results of the 5-fold cross-validation study aimed at validating MAP4C SVM, features SVM, RNN, and GPT-3.5 turbo performance for a) antimicrobial activity and b) hemolysis predictions. The mean performance across the 5 cross-validations for each metric is shown as a bar, the standard deviation is displayed with an error bar. The results confirmed earlier observations but showed notably higher performances for the RNN compared to the one-shot screening experiment. Both the RNN and MAP4C SVM demonstrated comparable performances.



The results, depicted in Fig. 1a for antimicrobial activity prediction and Fig. 1b for hemolysis prediction, confirmed our earlier observations (performances in Table S2†). Notably, the RNN performances were higher than those observed in the screening experiment, and were clearly above those of GTP-3.5. Furthermore, both the RNN and MAP4C SVM demonstrated comparable performances, indicating the validity of both approaches in predicting antimicrobial activity and hemolysis. The finding that simpler machine learning architectures, like SVM, can rival the performance of more complex RNNs in predicting antimicrobial activity and hemolysis is particularly interesting. A comparison with models trained on similar datasets, which achieve similar performances as reported in

this study, further reinforces the consistency of our findings.^{19–21}

This raises questions about the importance of model architecture *versus* foundational elements such as data quality and feature engineering. It suggests that a balanced approach, prioritizing optimization of these foundational components, could prove more beneficial than focusing solely on model complexity.

Data visualization

The high performance achieved by the SVM trained on the MAP4C fingerprint suggested that the nearest neighbor relationships in the MAP4C feature space could be sufficient

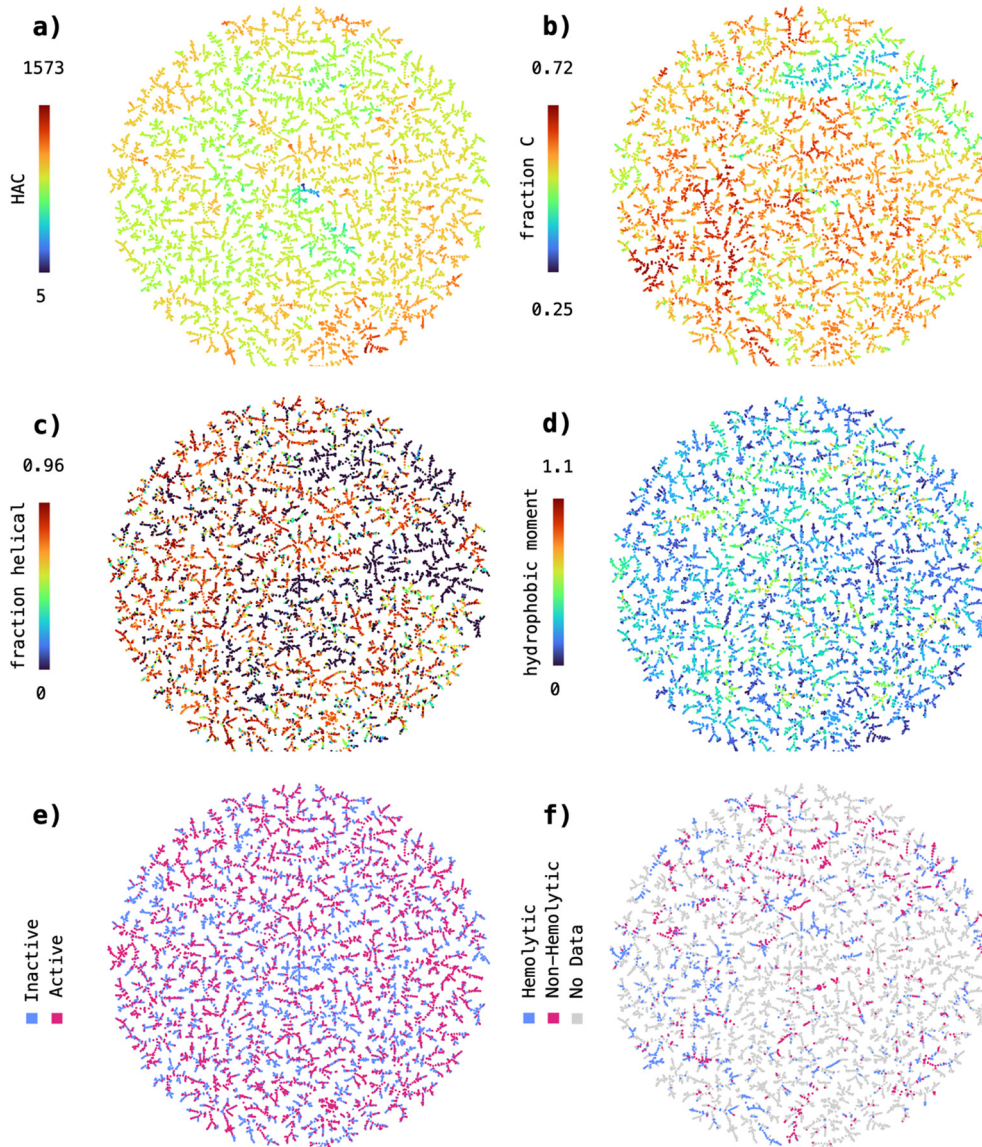


Fig. 2 Chemical space covered by the 9548 peptide sequences with annotated antimicrobial activity extracted from the Database of Antimicrobial Activity and Structure of Peptides (DBAASP). The sequences are encoded using the MAP4C fingerprint and the resulting 2048-dimensional space reduced to 2D using TMAP. The sequences in the 2D TMAP were colored based on a) heavy atom count, b) fraction of carbon atoms, c) predicted fraction of helical residues, d) hydrophobic moment, e) annotated antimicrobial activity and f) annotated hemolysis.

to distinguish active from inactive and hemolytic from non-hemolytic peptide sequences. In our previous work, we observed that the MAP4 fingerprint³⁸ correctly clustered natural products, taken from the COCONUT database,³⁹ according to their organism of origin.^{40,41} In analogy to our previous work, we were curious to see whether a spatial separation of actives/inactives and hemolytic/non-hemolytic sequences can be obtained from encoding with MAP4C, the chiral version of MAP4, possibly explaining the good performance of the MAP4C SVM model. For this, we reduced the 2048-dimensional feature space of MAP4C to 2D using the dimensionality reduction method TMAP,⁴² and used the obtained visualization to display a set of molecular properties.

First, we wanted to confirm that the TMAP visualization aligns with intuitive distributions of structural features relevant for peptides. For that, we colored the data points based on their heavy atom count (HAC), an indicator of molecular size, and fraction of carbon atoms (fraction C), a simple proxy for the hydrophobicity of a peptide sequence. The TMAP revealed visible clusters for both, HAC (Fig. 2a) and fraction C (Fig. 2b), indicating that the reduced MAP4C features can reliably represent simple molecular descriptors in the underlying chemical space.

Following this first observation, we wanted to test if we can detect clusters within TMAP visualizations of more complex physicochemical properties, such as the predicted fraction of helical residues (Fig. 2c) and the hydrophobic moment (Fig. 2d). In both cases, we could not detect large homogenous clusters as was the case for HAC and fraction C. However, the data formed a large number of small local clusters, indicating that the nearest neighbor relationships in the MAP4C feature space can possibly be used to distinguish sequences with high helicity/hydrophobicity opposed to sequences with low helicity/hydrophobicity.

Finally, we analysed the distribution of active *versus* inactive (Fig. 2e) and hemolytic *versus* non-hemolytic (Fig. 2f) sequences in the MAP4C chemical space. Similarly to the visualizations of predicted fraction of helical residues and hydrophobic moment, active and inactive or hemolytic and non-hemolytic sequences are spatially separated in a large number of small, local clusters. This finding is particularly interesting as it suggests that nearest neighbor relationships in the MAP4C feature space are sufficient to separate peptide sequences based on their antimicrobial activity and hemolysis. It further provides an explanation to the good performance obtained with the MAP4C SVM, which can leverage the nearest neighbor relationships stored in the MAP4C fingerprint feature space when provided with a custom Jaccard kernel function.

Conclusion

In the present study we investigated the potential of LLMs as predictive tools for antimicrobial activity and hemolysis of peptide sequences. We assessed that fine-tuning GPT models

in cloud is a relatively easy and fast process as access through the API eliminates the need to buy expensive hardware and requires little technical expertise. Duration of fine-tuning was short, and the associated costs were low (Table S3†). In contrast to cloud-based fine-tuning, local model training involves setting up and maintaining hardware, which can be costly and require technical expertise. While less complex models like RNNs and SVMs have lower hardware requirements, training larger models such as LLMs locally can pose challenges in terms of scalability, as one can rapidly face limitations in terms of hardware capacity and maintenance costs.

However, the lack of control over the training environment in cloud-based approaches raises concerns regarding reproducibility of scientific results. In the course of this study, we had originally fine-tuned GPT-3 models Ada, Babbage and Curie. These models performed slightly better than the reference models, even achieving good performances in low data regimes. Unfortunately, these models were later deprecated by OpenAI and their performance cannot be reproduced. When fine-tuning a newer iteration of GPT-3 (GPT-3.5 Turbo), we observed a significant decrease in performance for the same task. We attribute the drop in performance to the increasing optimization of LLMs for conversational interactions, which may negatively impact their effectiveness in out-of-scope predictive tasks. These findings highlight the potential risk of how not controlling one's own models can compromise the reproducibility and reliability of scientific results.

The aforementioned findings suggest a diminishing suitability of chat oriented LLMs for classification tasks over time, a function beyond their intended design. This observation specifically applies to LLMs tailored for conversational or human interaction purposes, rather than specialized LLMs trained on domain-specific data. Unfortunately, the latter do not provide the ease of access and usability that GPT models do. Consequently, we expect that LLMs will increasingly be employed in human interaction settings, facilitating the integration of various chemical tools through natural language interfaces as is being pioneered by Bran³¹ and Boiko *et al.*³²

Finally, we could demonstrate in the present study that classical machine learning techniques, such as SVMs trained on MAP4C fingerprint encodings, can achieve state-of-the-art performance in the prediction of antimicrobial activity and hemolysis. This finding is especially interesting, as it showcases that good performance can be achieved by less complex models, putting the emphasis on data quality rather than model complexity.

Code availability

The source codes and datasets used for this study are available at https://github.com/reymond-group/LLM_classifier.



Author contributions

MO designed and realized the project and wrote the paper. JLR designed and supervised the project and wrote the paper. Both authors read and approved the final manuscript.

Conflicts of interest

There is no conflict of interest to declare.

Acknowledgements

This work was supported by the Swiss National Science Foundation (200020_178998) and the European Research Council (885076). MO thanks Sacha Javor for the helpful discussion and comments.

References

- 1 M. Lakemeyer, W. Zhao, F. A. Mandl, P. Hammann and S. A. Sieber, Thinking Outside the Box—Novel Antibacterials To Tackle the Resistance Crisis, *Angew. Chem., Int. Ed.*, 2018, 57(44), 14440–14475, DOI: [10.1002/anie.201804971](#).
- 2 M. Magana, M. Pushpanathan, A. L. Santos, L. Leanse, M. Fernandez, A. Ioannidis, M. A. Giulianotti, Y. Apidianakis, S. Bradfute, A. L. Ferguson, A. Cherkasov, M. N. Seleem, C. Pinilla, C. De La Fuente-Nunez, T. Lazaridis, T. Dai, R. A. Houghten, R. E. W. Hancock and G. P. Tegos, The Value of Antimicrobial Peptides in the Age of Resistance, *Lancet Infect. Dis.*, 2020, 20(9), e216–e230, DOI: [10.1016/S1473-3099\(20\)30327-3](#).
- 3 N. Mookherjee, M. A. Anderson, H. P. Haagsman and D. J. Davidson, Antimicrobial Host Defence Peptides: Functions and Clinical Potential, *Nat. Rev. Drug Discovery*, 2020, 19(5), 311–332, DOI: [10.1038/s41573-019-0058-8](#).
- 4 M. D. T. Torres, S. Sothiselvam, T. K. Lu and C. De La Fuente-Nunez, Peptide Design Principles for Antimicrobial Applications, *J. Mol. Biol.*, 2019, 431(18), 3547–3567, DOI: [10.1016/j.jmb.2018.12.015](#).
- 5 A. Capecchi and J.-L. Reymond, Peptides in Chemical Space, *Med. Drug Discovery*, 2021, 9, 100081, DOI: [10.1016/j.medidd.2021.100081](#).
- 6 A. T. Müller, J. A. Hiss and G. Schneider, Recurrent Neural Network Model for Constructive Peptide Design, *J. Chem. Inf. Model.*, 2018, 58(2), 472–479, DOI: [10.1021/acs.jcim.7b00414](#).
- 7 D. Veltri, U. Kamath and A. Shehu, Deep Learning Improves Antimicrobial Peptide Recognition, *Bioinformatics*, 2018, 34(16), 2740–2747, DOI: [10.1093/bioinformatics/bty179](#).
- 8 S. Liu, Novel 3D Structure Based Model for Activity Prediction and Design of Antimicrobial Peptides, *Sci. Rep.*, 2018, 8, 11189, DOI: [10.1038/s41598-018-29566-5](#).
- 9 X. Su, J. Xu, Y. Yin, X. Quan and H. Zhang, Antimicrobial Peptide Identification Using Multi-Scale Convolutional Network, *BMC Bioinf.*, 2019, 20(1), 730, DOI: [10.1186/s12859-019-3327-y](#).
- 10 B. Vishnepolsky, G. Zaalishvili, M. Karapetian, T. Nasrashvili, N. Kuljanishvili, A. Gabrielian, A. Rosenthal, D. E. Hurt, M. Tartakovsky, M. Grigolava and M. Pirtskhalava, De Novo Design and In Vitro Testing of Antimicrobial Peptides against Gram-Negative Bacteria, *Pharmaceuticals*, 2019, 12(2), 82, DOI: [10.3390/ph12020082](#).
- 11 F. Plisson, O. Ramírez-Sánchez and C. Martínez-Hernández, Machine Learning-Guided Discovery and Design of Non-Hemolytic Peptides, *Sci. Rep.*, 2020, 10(1), 16581, DOI: [10.1038/s41598-020-73644-6](#).
- 12 J. Yan, P. Bhadra, A. Li, P. Sethiya, L. Qin, H. K. Tai, K. H. Wong and S. W. I. Siu, Deep-AmPEP30: Improve Short Antimicrobial Peptides Prediction with Deep Learning, *Mol. Ther.–Nucleic Acids*, 2020, 20, 882–894, DOI: [10.1016/j.omtn.2020.05.006](#).
- 13 A. Capecchi, X. Cai, H. Personne, T. Köhler, C. van Delden and J.-L. Reymond, Machine Learning Designs Non-Hemolytic Antimicrobial Peptides, *Chem. Sci.*, 2021, 12(26), 9221–9232, DOI: [10.1039/D1SC01713F](#).
- 14 E. Zakharova, M. Orsi, A. Capecchi and J. Reymond, Machine Learning Guided Discovery of Non-Hemolytic Membrane Disruptive Anticancer Peptides, *ChemMedChem*, 2022, 17(17), DOI: [10.1002/cmdc.202200291](#).
- 15 G. Liu, D. B. Catacutan, K. Rathod, K. Swanson, W. Jin, J. C. Mohammed, A. Chiappino-Pepe, S. A. Syed, M. Frangis, K. Rachwalski, J. Magolan, M. G. Surette, B. K. Coombes, T. Jaakkola, R. Barzilay, J. J. Collins and J. M. Stokes, Deep Learning-Guided Discovery of an Antibiotic Targeting *Acinetobacter Baumannii*, *Nat. Chem. Biol.*, 2023, 19, 1342–1350, DOI: [10.1038/s41589-023-01349-8](#).
- 16 F. Wan and C. De La Fuente-Nunez, Mining for Antimicrobial Peptides in Sequence Space, *Nat. Biomed. Eng.*, 2023, 7, 707–708, DOI: [10.1038/s41551-023-01027-z](#).
- 17 M. D. C. Aguilera-Puga and F. Plisson, Structure-Aware Machine Learning Strategies for Antimicrobial Peptide Discovery, *Research Square*, 2024, preprint, DOI: [10.21203/rs.3.rs-3938402/v1](#).
- 18 F. Wan, F. Wong, J. J. Collins and C. De La Fuente-Nunez, Machine Learning for Antimicrobial Peptide Identification and Design, *Nat. Rev. Bioeng.*, 2024, DOI: [10.1038/s44222-024-00152-x](#).
- 19 P. B. Timmons and C. M. Hewage, HAPPENN Is a Novel Tool for Hemolytic Activity Prediction for Therapeutic Peptides Which Employs Neural Networks, *Sci. Rep.*, 2020, 10(1), 10869, DOI: [10.1038/s41598-020-67701-3](#).
- 20 M. M. Hasan, N. Schaduagrat, S. Basith, G. Lee, W. Shoombuatong and B. Manavalan, HLPpred-Fuse: Improved and Robust Prediction of Hemolytic Peptide and Its Activity by Fusing Multiple Feature Representation, *Bioinformatics*, 2020, 36(11), 3350–3356, DOI: [10.1093/bioinformatics/btaa160](#).
- 21 M. Ansari and A. D. White, Serverless Prediction of Peptide Properties with Recurrent Neural Networks, *J. Chem. Inf. Model.*, 2023, 63(8), 2546–2553, DOI: [10.1021/acs.jcim.2c01317](#).
- 22 S. Hochreiter and J. Schmidhuber, Long Short-Term Memory, *Neural Comput.*, 1997, 9(8), 1735–1780, DOI: [10.1162/neco.1997.9.8.1735](#).
- 23 K. Cho, B. van Merriënboer, D. Bahdanau and Y. Bengio, On the Properties of Neural Machine Translation: Encoder-



- Decoder Approaches, *arXiv*, 2014, preprint, DOI: [10.48550/arXiv.1409.1259](https://doi.org/10.48550/arXiv.1409.1259), (accessed 2023-05-31).
- 24 A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, L. Kaiser and I. Polosukhin, Attention Is All You Need, *arXiv*, 2017, preprint, DOI: [10.48550/arXiv.1706.03762](https://doi.org/10.48550/arXiv.1706.03762), (accessed 2023-05-31).
 - 25 T. B. Brown, B. Mann, N. Ryder, M. Subbiah, J. Kaplan, P. Dhariwal, A. Neelakantan, P. Shyam, G. Sastry, A. Askell, S. Agarwal, A. Herbert-Voss, G. Krueger, T. Henighan, R. Child, A. Ramesh, D. M. Ziegler, J. Wu, C. Winter, C. Hesse, M. Chen, E. Sigler, M. Litwin, S. Gray, B. Chess, J. Clark, C. Berner, S. McCandlish, A. Radford, I. Sutskever and D. Amodei, Language Models Are Few-Shot Learners, *arXiv*, 2020, preprint, DOI: [10.48550/arXiv.2005.14165](https://doi.org/10.48550/arXiv.2005.14165), (accessed 2023-05-31).
 - 26 K. M. Jablonka, P. Schwaller, A. Ortega-Guerrero and B. Smit, Leveraging Large Language Models for Predictive Chemistry, *Nat. Mach. Intell.*, 2024, **6**(2), 161–169, DOI: [10.1038/s42256-023-00788-1](https://doi.org/10.1038/s42256-023-00788-1).
 - 27 A. M. Bran and P. Schwaller, Transformers and Large Language Models for Chemistry and Drug Discovery, *arXiv*, 2023, preprint, DOI: [10.48550/arXiv.2310.06083](https://doi.org/10.48550/arXiv.2310.06083).
 - 28 T. Guo, K. Guo, B. Nan, Z. Liang, Z. Guo, N. V. Chawla, O. Wiest and X. Zhang, What Can Large Language Models Do in Chemistry? A Comprehensive Benchmark on Eight Tasks, part of Advances in Neural Information Processing Systems, *NeurIPS Proceedings*, 2023, vol. 36, pp. 59662–59688.
 - 29 C. M. Castro Nascimento and A. S. Pimentel, Do Large Language Models Understand Chemistry? A Conversation with ChatGPT, *J. Chem. Inf. Model.*, 2023, **63**(6), 1649–1655, DOI: [10.1021/acs.jcim.3c00285](https://doi.org/10.1021/acs.jcim.3c00285).
 - 30 A. D. White, G. M. Hocky, H. A. Gandhi, M. Ansari, S. Cox, G. P. Wellawatte, S. Sasmal, Z. Yang, K. Liu, Y. Singh and W. J. Peña Ccoa, Assessment of Chemistry Knowledge in Large Language Models That Generate Code, *Digital Discovery*, 2023, **2**(2), 368–376, DOI: [10.1039/D2DD00087C](https://doi.org/10.1039/D2DD00087C).
 - 31 A. M. Bran, S. Cox, A. D. White and P. Schwaller, ChemCrow: Augmenting Large-Language Models with Chemistry Tools, *arXiv*, 2023, preprint, DOI: [10.48550/arXiv.2304.05376](https://doi.org/10.48550/arXiv.2304.05376), (accessed 2023-05-31).
 - 32 D. A. Boiko, R. MacKnight, B. Kline and G. Gomes, Autonomous Chemical Research with Large Language Models, *Nature*, 2023, **624**(7992), 570–578, DOI: [10.1038/s41586-023-06792-0](https://doi.org/10.1038/s41586-023-06792-0).
 - 33 K. M. Jablonka, Q. Ai, A. Al-Feghali, S. Badhwar, J. D. Bocarsly, A. M. Bran, S. Bringuier, L. C. Brinson, K. Choudhary, D. Circi, S. Cox, W. A. De Jong, M. L. Evans, N. Gastellu, J. Genzling, M. V. Gil, A. K. Gupta, Z. Hong, A. Imran, S. Kruschwitz, A. Labarre, J. Lála, T. Liu, S. Ma, S. Majumdar, G. W. Merz, N. Moitessier, E. Moubarak, B. Mouriño, B. Pelkie, M. Pieler, M. C. Ramos, B. Ranković, S. G. Rodrigues, J. N. Sanders, P. Schwaller, M. Schwarting, J. Shi, B. Smit, B. E. Smith, J. Van Herck, C. Völker, L. Ward, S. Warren, B. Weiser, S. Zhang, X. Zhang, G. A. Zia, A. Scourtas, K. J. Schmidt, I. Foster, A. D. White and B. Blaiszik, 14 Examples of How LLMs Can Transform Materials Science and Chemistry: A Reflection on a Large Language Model Hackathon, *Digital Discovery*, 2023, **2**(5), 1233–1250, DOI: [10.1039/D3DD000113J](https://doi.org/10.1039/D3DD000113J).
 - 34 M. Orsi and J.-L. Reymond, One Chiral Fingerprint to Find Them All, *ChemRxiv*, 2023, preprint, DOI: [10.26434/chemrxiv-2023-33j02](https://doi.org/10.26434/chemrxiv-2023-33j02).
 - 35 G. Gogoladze, M. Grigolava, B. Vishnepolsky, M. Chubinidze, P. Duroux, M.-P. Lefranc and M. Pirtskhalava, DBAASP: Database of Antimicrobial Activity and Structure of Peptides, *FEMS Microbiol. Lett.*, 2014, **357**(1), 63–68, DOI: [10.1111/1574-6968.12489](https://doi.org/10.1111/1574-6968.12489).
 - 36 R. Heffernan, K. Paliwal, J. Lyons, J. Singh, Y. Yang and Y. Zhou, Single-sequence-based Prediction of Protein Secondary Structures and Solvent Accessibility by Deep Whole-sequence Learning, *J. Comput. Chem.*, 2018, **39**(26), 2210–2216, DOI: [10.1002/jcc.25534](https://doi.org/10.1002/jcc.25534).
 - 37 D. Eisenberg, R. M. Weiss and T. C. Terwilliger, The Helical Hydrophobic Moment: A Measure of the Amphiphilicity of a Helix, *Nature*, 1982, **299**(5881), 371–374, DOI: [10.1038/299371a0](https://doi.org/10.1038/299371a0).
 - 38 A. Capecchi, D. Probst and J.-L. Reymond, One Molecular Fingerprint to Rule Them All: Drugs, Biomolecules, and the Metabolome, *Aust. J. Chem.*, 2020, **12**(1), 43, DOI: [10.1186/s13321-020-00445-4](https://doi.org/10.1186/s13321-020-00445-4).
 - 39 M. Sorokina, P. Merseburger, K. Rajan, M. A. Yirik and C. Steinbeck, COCONUT Online: Collection of Open Natural Products Database, *Aust. J. Chem.*, 2021, **13**(1), 2, DOI: [10.1186/s13321-020-00478-9](https://doi.org/10.1186/s13321-020-00478-9).
 - 40 A. Capecchi and J.-L. Reymond, Assigning the Origin of Microbial Natural Products by Chemical Space Map and Machine Learning, *Biomolecules*, 2020, **10**(10), 1385, DOI: [10.3390/biom10101385](https://doi.org/10.3390/biom10101385).
 - 41 A. Capecchi and J.-L. Reymond, Classifying Natural Products from Plants, Fungi or Bacteria Using the COCONUT Database and Machine Learning, *Aust. J. Chem.*, 2021, **13**(1), 82, DOI: [10.1186/s13321-021-00559-3](https://doi.org/10.1186/s13321-021-00559-3).
 - 42 D. Probst and J.-L. Reymond, Visualization of Very Large High-Dimensional Data Sets as Minimum Spanning Trees, *Aust. J. Chem.*, 2020, **12**(1), 12, DOI: [10.1186/s13321-020-0416-x](https://doi.org/10.1186/s13321-020-0416-x).

