

<https://doi.org/10.1038/s41524-025-01687-2>

Machine learning on multiple topological materials datasets



Yuqing He^{1,2}, Pierre-Paul De Breuck¹, Hongming Weng^{1,2}✉, Matteo Giantomassi¹ & Gian-Marco Rignanese^{1,3,4}✉

A dataset of 35,608 materials with their topological properties is constructed by combining the density functional theory (DFT) results of Materials and the Topological Materials Database. Thanks to this, machine-learning approaches are developed to categorize materials into five distinct topological types, with the XGBoost model achieving an impressive 85.2% classification accuracy. By conducting generalization tests on different sub-datasets, differences are identified between the original datasets in terms of topological types, chemical elements, unknown magnetic compounds, and feature space coverage. Their impact on model performance is analyzed. Turning to the simpler binary classification between trivial insulators and nontrivial topological materials, three different approaches are also tested. Key characteristics influencing material topology are identified, with the maximum packing efficiency and the fraction of p valence electrons being highlighted as critical features.

In the past decade, topological electronic materials have drawn considerable attention due to their importance for both fundamental science and next-generation technological applications^{1–3}. These materials exhibit unique topological configurations in their electronic band structures, resulting in peculiar electronic properties^{2,3}. A central and long-standing question in this field is how to determine whether a given material is topologically trivial or not. Thanks to advancements in symmetry indicator theory⁴ and topological quantum chemistry theory⁵ along with the use of first-principles calculations^{6–8}, it is now possible to categorize topologically non-trivial materials (NTMs) according to their symmetry indicators or elementary band representations (EBRs) and compatibility relations (CR)^{9–11}. Among these NTMs, topological semimetals (TSMs) represent a key class characterized by electronic bands that intersect at discrete points or along lines in momentum space. Their irreducible representations violate CR, resulting in gapless states near the Fermi level. These topological nodes cannot be removed by symmetry-preserving perturbations. Depending on the position of these nodes identified according to the symmetry representations of the corresponding bands, TSMs can be further grouped into high-symmetry-point semimetals (HSPSMs) and high-symmetry-line semimetals (HLSMs)⁶. Depending on the symmetry indicators, one then identifies the topological insulators (TIs) and the topological crystalline insulators (TCIs), which have a set of valence bands that satisfy the CR but cannot be decomposed into linear combination of EBRs⁵. Using high-throughput calculations integrating density functional theory (DFT)^{12,13} and topological quantum chemistry theory or symmetry indicator theory, tens

of thousands of topological materials were detected by analyzing the symmetry of the wavefunctions of crystalline compounds from the Inorganic Crystal Structure Database¹⁴ (ICSD), leading to the compilation of several databases^{6–8}. Soon after, these two theories based on symmetry representations of wavefunctions have been extended to magnetic ordered materials^{15–17}.

However, from an experimental standpoint, the main characteristics of materials that may affect their topological properties, especially those which provide helpful chemical insights, are still unclear, hindering the design of new topological materials. Machine learning (ML)¹⁸ techniques offer a novel approach to address these shortcomings. By exploring existing data, they can potentially identify the features that are critical for obtaining topological properties without the need to resort to heavy computations. In this framework, several studies have already been conducted. For example, neural networks and k-means clustering have been used to learn from the Hamiltonian^{19–22}. Compressed-sensing²³, gradient boosted trees (GBTs)²⁴, and other methods have also been employed to learn from ab-initio data^{25–31}. Claussen et al.²⁶ trained a GBT model with 35,009 symmetry-based entries from the Topological Materials Database (TMD)^{7,32}. By testing the effect of different features, their model reached an accuracy of 87.0% for classifying materials into five subclasses: Trivial Insulator (TrIs), two types of TSMs (Enforced Semimetals (ESs) or Enforced Semimetals with Fermi Degeneracy (ESFDs)), and two types of TIs (Split EBRs (SEBRs) or Not a Linear Combination of EBRs (NLC)). Andrejevic et al.³⁰ utilized 16,458 inorganic materials from TMD to develop a neural network classifier capable of

¹Institute of Condensed Matter and Nanosciences, UCLouvain, Louvain-la-Neuve, Belgium. ²Beijing National Laboratory for Condensed Matter Physics and Institute of Physics, Chinese Academy of Sciences, Beijing, China. ³WEL Research Institute, Wavre, Belgium. ⁴School of Materials Science and Engineering, Northwestern Polytechnical University, Xi'an, Shaanxi, China. ✉e-mail: hmweng@iphy.ac.cn; gian-marco.rignanese@uclouvain.be

distinguishing NTMs and TrIs based on X-ray absorption near-edge structure (XANES) spectra. It achieved a F_1 score of 89% for predicting NTMs from their XANES signatures. In 2023, Ma et al.³¹ introduced a parameter named *topogivity* for each chemical element measuring its tendency to form topological materials. They employed support vector machine (SVM)³³ to learn the topogivities of elements based on a dataset of 9026 materials from Tang et al.⁸. This approach achieved an average accuracy of 82.7% in an 11×10 -fold NCV procedure. Claussen et al.²⁶ found that the topology does not depend much on the particular positions of atoms in the crystal lattice. This is in contrast with the previous two findings indicating that the local environment of atoms in compounds is a decisive factor since both XANES and topogivity are sensitive to the elements and their local chemical environment. In addition, these previous studies used mainly generic ML algorithms, overlooking newly designed ones specifically tailored for materials science and which have demonstrated excellent performance^{34–37}. Finally, we would like to point out that the outcomes of these studies are difficult to compare due to the diverse range of crystal systems considered (including specific classes of systems, 2D materials, or bulk 3D materials) and the variations in utilized databases, material classes, and types of features incorporated in the ML model construction.

In this paper, we conduct data curation to compile a comprehensive dataset consisting of 35,608 entries from *Materiae*⁶ and TMD. This new dataset is then used to train models for classifying a material into five types as TrI, or NTM, specifically HSPSM, HSLSM, TI, and TCI. Using nested cross-validation (NCV)³⁸ on the data from *Materiae*, we first benchmark five different approaches: namely Random Forest (RF)³⁹, XGBoost⁴⁰ (an implementation of GBTs), Automatminer (AMM)³⁵, the Material Optimal Descriptor Network (MODNet)^{36,41}, and the Materials Graph Network (MEGNet)³⁴. We find that XGBoost performs the best. We then test this model on the complete dataset, achieving a mean NCV accuracy of 82.9% for the classification into the five types. We discuss the NCV results and the train-test procedures, highlighting differences between the datasets that affect the scores. Finally, we compare XGBoost with topogivity³¹ and t-distributed Stochastic Neighbor Embedding (t-SNE)⁴² for the binary classification between TrIs and NTMs. XGBoost performs the best with 92.4% accuracy. Additionally, we investigate the key factors influencing the topology of materials. Our analysis reveals that the maximum packing efficiency (MPE) and the fraction of p valence electrons (FPV) are the factors that contribute the most to the distinction between TrIs and NTMs. It is noted that MPE is a structure-based feature that represents the maximum packing efficiency of atoms within a crystal lattice and FPV is a composition-based feature that indicates the fraction of p electrons versus all valence electrons. We think this finding is reasonable and these features could be used as a heuristic for exploring new topological materials.

Results

Data curation

Two datasets are constructed in this work, named M and T . The dataset M is extracted from *Materiae* following a thorough data curation procedure as described below, resulting in 25,683 compounds. Similarly, the dataset T is constructed from the Topological Materials Database, resulting in 24,156 compounds after cleaning. It should be noted that the names of the topological types in these two databases are different. Therefore, we here establish a correspondence between them during the curation process. Figure 1 depicts the type distribution for the two datasets, their intersection ($M \cap T$), and differences ($M \setminus T$ and $T \setminus M$). Roughly the same distribution is found, with a majority of TrIs and around 30% NTMs.

The data curation proceeds as follows. For the dataset M , we initially query *Materiae*, which includes 26,120 materials that are neither magnetic materials (i.e., for which the magnetic moment would be higher than 0.1 Bohr Magnetons per unit cell according to the Materials Project⁴³ (MP) record) nor conventional metals (i.e., systems with an odd number of electrons per unit cell). By keeping only the results that were computed including spin-orbit coupling, an initial dataset of 25,895 materials (named MAT) is obtained including their topological properties. Subsequently, the

dataset M is constructed by removing those materials with labels conflicting with the dataset T (see last paragraph). All the records in MAT are indexed by their unique MP-ID.

For the dataset T , we start from the data available in the Topological Materials Database³², which includes 73,234 compounds indexed by their ICSD-ID and grouped into 38,298 unique materials by common chemical formula, space group, and topological properties as determined from their calculated electronic structure. As some of the pre-assigned MP-IDs were found to be wrong, we decided to control them systematically with the *structure_matcher* of PYMATGEN⁴⁴ (using its default tolerance settings) and make our own MP-ID assignment. Given a set of compounds grouped as one unique material, we distinguish three cases to assign the MP-ID and the corresponding structure. First, when none of the compounds has an assigned MP-ID, one structure of the set is randomly selected and the corresponding MP-ID is indicated as not available. Second, when the compounds are associated with at most one MP-ID, one structure of the set is again randomly selected. We then check whether it matches the MP structure corresponding to the indicated MP-ID. If it does, the MP-ID is assigned to the structure. If not, the MP-ID is indicated as not available. Finally, when more than one MP-ID appears in the set, these MP-IDs are first ranked according to their energy above hull. Then, the different structures in the set are compared with the MP structures corresponding to these MP-IDs starting from the lowest in energy. In case of match, the corresponding MP-ID (i.e., the one with the lowest energy) is assigned to the structure. In case of absence of match with any of the ranked MP-IDs, the MP-ID is indicated as not available. At the end of the process, only one of the structures associated with the same MP-IDs is kept. The compounds are then sorted adopting the same classification as in *Materiae*. First, they undergoes the same curation as the one described for the dataset M : excluding magnetic materials and conventional metals. The materials containing rare-earth elements (Pr, Nd, Pm, Sm, Tb, Dy, Ho, Er, Tm, Yb, Lu, Sc) are also removed. This is done because, for these elements, the results of *Materiae* and the Topological Materials Database were obtained from calculations performed using pseudopotentials with a different number of valence electrons (typically odd in one case and even in the other). Furthermore, we label the resulting data according to *Materiae*'s definition. For TSMs, the mapping is rather simple: ESFDs correspond to HSPSMs and ESs to HSLSMs. In contrasts, for TIs, the mapping is more complex. We classify SEBRs and NLCs as TIs or TCIs as follows. The materials in the spacegroups 174, 187, 189, 188, or 190 are all labeled as TCIs. The others are labeled according to the parity of the last topological indices, odd ones as TIs while even ones as TCIs. The next curation step consists in removing the materials with duplicate MP-IDs, as well the 673 compounds with the same MP-ID but conflicting topological types. At the end, we were left with a total of 24,368 items with an assigned MP-ID (sometimes indicated as not available) and sorted according to the same classification as *Materiae*. Thanks to the curation performed, the compounds in the two datasets can easily be related based on their assigned MP-IDs. On this basis, we further removed 212 materials present in both datasets but with differing types, leaving 24,156 compounds in T .

At the end of the construction of the datasets T and M , our global dataset $M \cup T$ contains a total of 35,608 materials while the intersection $M \cap T$ consists of 14,231 compounds, as shown in Fig. 1.

Model

In order to select a model for further training and analysis, we first perform a benchmark on the MAT dataset. Five different models are used: two generic ML algorithms (RF and XGBoost), and three well-developed algorithms in the field of material science (AMM, MODNet, and MEGNet). Moreover, for each method, two procedures are considered for the multiclass classification: either a direct multiclass classification (which gives the 5 possible labels as output) or a hierarchical binary classification (multiple models are trained following a tree such that each leaf represents a class). Figure 8 schematically represents these two procedures, with their respective accuracies. The highest accuracy (85.2%) is obtained with XGBoost using the direct

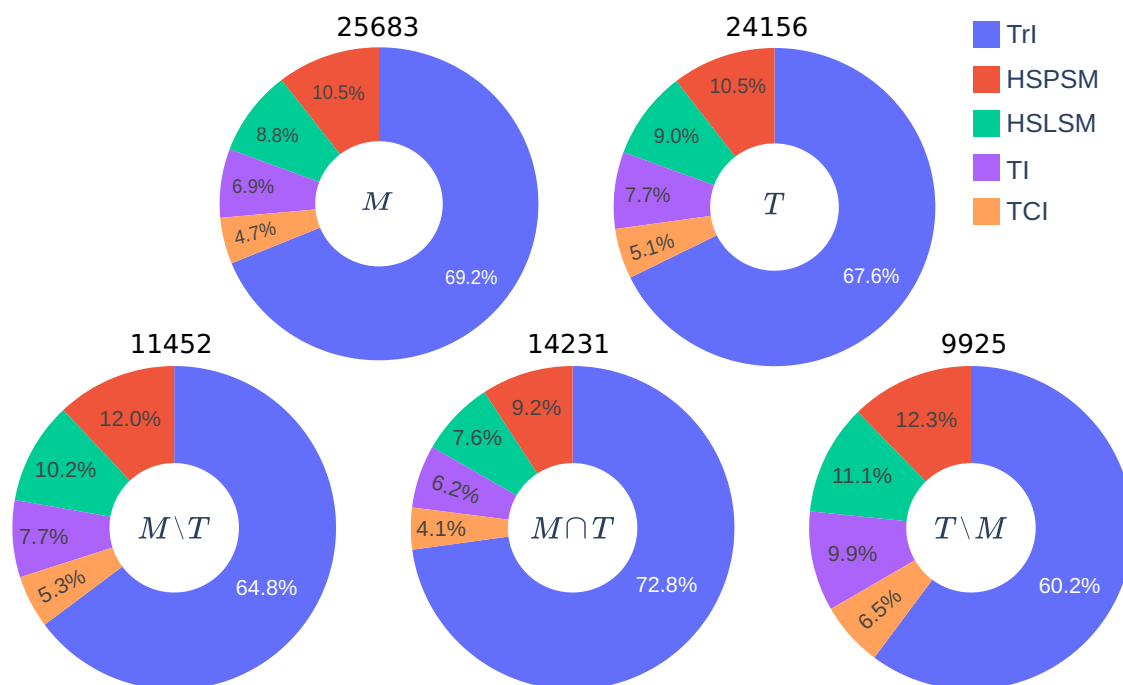
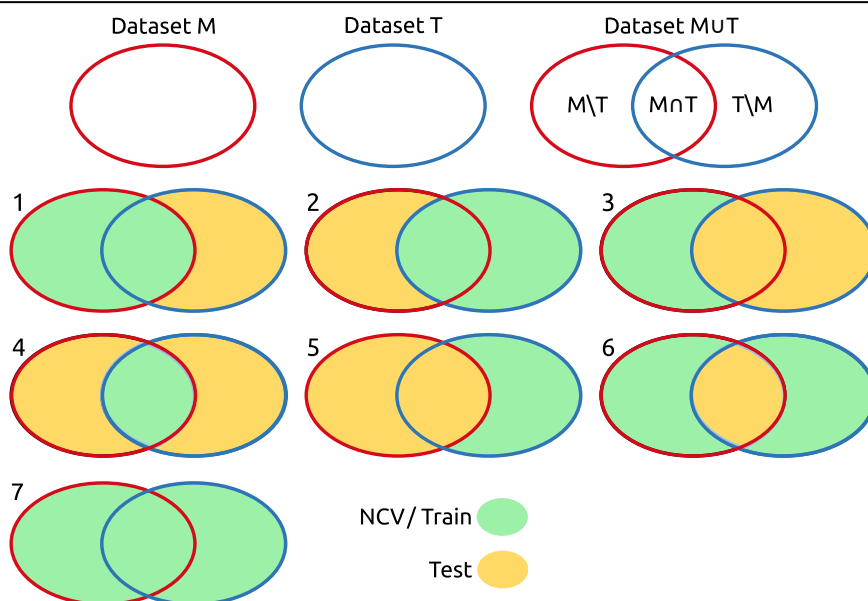


Fig. 1 | Composition of the datasets. Composition in terms of the different topological types of the datasets M (constructed from Materiae) and T (originating from the Topological Materials Database) as well as of their differences ($M \setminus T$ and $T \setminus M$) and intersection ($M \cap T$).

Fig. 2 | Diagram of the seven generalization tests.

The datasets M and T are circled in red and blue, respectively. The union dataset $M \cup T$ can be split into three parts: $M \setminus T$, $M \cap T$, and $T \setminus M$. In each test, the dataset used for the training and the NCV of the ML model is filled in green, while the dataset used for testing it is filled in yellow.



multiclass classification. It is therefore used in the remainder of the work on all the data. More details about the benchmark are provided in Section “ML Models”, which discusses the feature engineering and hyperparameter configurations for each algorithm, along with the training process.

Generalization tests

In principle, the NCV score should provide a comprehensive assessment of how the training model performs on new data. However, the model trained on the dataset M , which shows excellent performance (with a NCV accuracy of 85.7%), is found not to generalize well on the dataset $T \setminus M$ leading to an accuracy of only 71.8% (i.e., a decrease of 13.9%). Therefore, in order to further investigate the model performance, we perform a series of generalization tests by training on the different datasets at our disposal: M , T ,

their differences $M \setminus T$ and $T \setminus M$, their intersection $M \cap T$, their union $M \cup T$ as well as the union of their differences ($M \setminus T \cup T \setminus M$).

The seven tests are schematically represented in Fig. 2, where the datasets M and T are circled in red and blue, respectively. In each test, a ML model is first trained on the training set, depicted in green. A five-fold NCV test is performed on the same data, followed by a generalization test on the test set depicted in yellow. The classification accuracy results obtained for each test are reported in Table 1, indicating the score obtained for the NCV, as well as on $M \setminus T$, $M \cap T$, $T \setminus M$, and $M \cup T$. Complementary metrics (i.e., F_1 score, precision, and recall) are reported in Table S1 in the Supplementary Information.

As discussed below, the previously mentioned generalization issue is still present. For the generalization tests performed with M , T , $M \setminus T$, and

Table 1 | Accuracy (in %) of the different nested cross-validation (NCV) and generalization tests depending on the training dataset

Test \ Train	M	T	$M \setminus T$	$M \cap T$	$T \setminus M$	$(M \setminus T) \cup (T \setminus M)$	$M \cup T$
NCV	85.7	80.7	84.1	85.1	72.1	79.9	82.9
$M \setminus T$	84.8*	80.3	84.1*	78.9	77.1	85.6*	85.8*
$M \cap T$	86.1*	86.0*	82.0	85.1*	81.6	83.6	86.6*
$T \setminus M$	71.8	73.2*	69.8	70.2	72.1*	73.3*	74.3*
$M \cup T$	81.8	80.6	79.3	79.0	77.5	81.4	82.9

The results on (part of) the training dataset evaluated by NCV are indicated by a star.

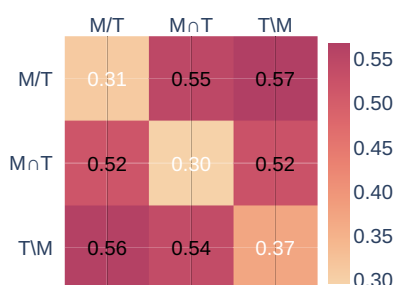


Fig. 3 | Heterogeneity metric between datasets: $M \setminus T$, $M \cap T$ and $T \setminus M$. The heterogeneity from dataset A (y-axis) to dataset B (x-axis) was computed as the average distance from each point in A to its 5-nearest neighbors in B, using a feature space defined by the top 47 Matminer features (44 continuous, 3 discrete). These features were selected based on their importance in training XGBoost models. Larger distances indicate higher dissimilarity, revealing the compositional differences between the datasets.

$M \cap T$ as the training set (in green), the NCV accuracy is significantly higher than the test accuracy (i.e., for the corresponding datasets in yellow). When training on the dataset M , the NCV accuracy on the sub-dataset $M \setminus T$ (84.8%) and $M \cap T$ (86.1%) is also much larger than the test accuracy (71.8% for $T \setminus M$). When training on the dataset T , the results are more nuanced with the NCV on the sub-dataset $M \cap T$ (86.0%) being higher than the test accuracy (80.3% for $M \setminus T$). But that on the sub-dataset $T \setminus M$ (73.2%) is not.

It is worth noting that, when training on $T \setminus M$ and $(M \setminus T) \cup (T \setminus M)$, the NCV accuracy (72.1% and 79.9%) is smaller than all test accuracy values (81.6% and 83.6% for $M \cap T$, respectively; as well as 77.1% for $M \setminus T$ in the former case). Finally, the accuracy on $T \setminus M$ is the lowest one whatever the training set.

All the other metrics (F_1 score, precision, and recall) reported in Table S1 show the same trend. All these observations indicate that predicting the topological type on the materials of the dataset $T \setminus M$ seems to be more difficult than on those of the dataset M (or its sub-datasets $M \setminus T$ and $M \cap T$). We propose four possible explanations for this bias (which are most probably combined).

The first reason is related to the distribution of the topological types in the datasets. As can be seen in Fig. 1, the proportion of TrIs is the lowest in $T \setminus M$, and the binary classification between TrIs and NTMs is much more accurate than the subsequent refined classifications of NTMs (see Fig. 8, Node 1 with respect to all the other nodes). Therefore, the proportion of TrIs affects the global accuracy.

The second rationalization is based on the distribution of the chemical elements in the datasets. Indeed, the accuracy of the model can be very low on compounds containing certain elements (e.g., as low as 37% on average for Gd), as illustrated in the Supplementary Information (Fig. S1). In particular, the following elements with a low average accuracy are more present in $T \setminus M$ than in any other dataset: Ne, Mn, Fe, Eu, Gd, Po, Rn, Ra, Am. To test

how this affects the global accuracy in each dataset, we recalculate the performance of the model when these elements are excluded. The corresponding accuracy, F_1 score, precision, and recall as well as the proportion of these materials are reported in the Supplementary Information (Table S2). In general, the performance is smaller when including the elements above. This decrease is more important for the dataset $T \setminus M$ (3% compared to 0.5% for the other datasets). This could be expected as it contains a larger fraction of the elements above.

Following upon this observation, we search for possibly problematic elements in the dataset $M \cup T$. Their detection is based on more quantitative criteria. First, the number of materials containing such problematic element should be larger than 30, for statistical reasons. Second, the accuracy for the compounds containing this element should be lower than 75%. Finally, the recall for those materials should be lower than the one for those without that element. Applying these criteria, the following elements are identified: Cr, Mn, Fe, Cu, Tc, Eu, Os, Np. Table S3 contains the accuracies, F_1 score, precision and recall based on the presence of the previous elements.

A third potential cause of the bias for the dataset $T \setminus M$ is that about half of its compounds have an unknown magnetic type, since they could not be assigned an MP-ID. Table S4 investigates both the impact of elements and the presence of magnetic information. As can be seen, excluding the selected elements in the datasets $M \setminus T$, $M \cap T$ or $T \setminus M$ improves the accuracy by 5.3%. Excluding compounds with missing magnetic information further improves the score by 1.2%.

To analyze the cumulative effect of the above three explanations, we define the datasets $\widetilde{M \setminus T}$, $\widetilde{M \cap T}$, and $\widetilde{T \setminus M}$. These are formed by selecting the same number of compounds (3372) in each original dataset ($M \setminus T$, $M \cap T$, and $T \setminus M$) adopting the same criteria as in Table S4 and in such a way that the distribution among the five different types is exactly the same (i.e., 2339 TrIs, 315 HSPSMs, 279 HSLSMs, 271 TIs, and 168 TCIs). As can be seen in the NCV results reported in Table S5, the accuracy in the three datasets (79.2%, 79.9%, and 77.3%) is much more similar (the largest difference decreased to 2.6% from the previous 13.9%).

Finally, a fourth possible reason is related to the coverage of the feature space by the datasets. The ML model performance on a given test set obviously depends on how close its points are from those of the training dataset (interpolative predictions are better than extrapolative ones). To evaluate this effect, a heterogeneity metric is used, as explained in detail in “Methods” (see Eq. (5)). It quantifies the similarity between the different datasets ($M \setminus T$, $M \cap T$, and $T \setminus M$), with a small heterogeneity leading in principle to a higher performance. The heterogeneity within each dataset (the diagonal part in Fig. 3) provides a reference value. Note that the heterogeneity in the dataset $T \setminus M$ is about 20% larger than the others. This may explain the trend in the NCV accuracy for the models trained on the datasets $\widetilde{M \setminus T}$, $\widetilde{M \cap T}$, and $\widetilde{T \setminus M}$: as expected the lower the heterogeneity, the higher the NCV score. Furthermore, the heterogeneity increases significantly in the off-diagonal elements. This explains why a model trained on a given dataset tends not to generalize well to the other datasets.

Binary classification

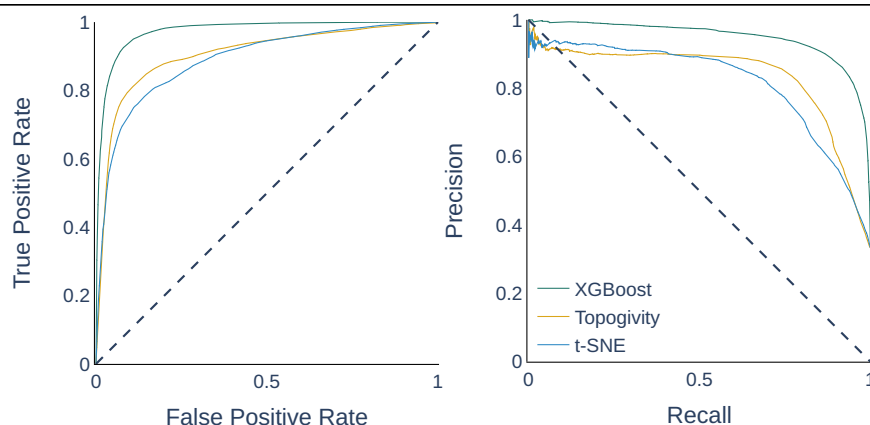
In order to try to identify the main factors that influence the topology of a material, we turn to the binary classification between TrIs and NTMs on the whole dataset $M \cup T$. NTMs are considered as positive and TrIs as negative. Thus, the precision measures the reliability of NTM predictions, and the recall measures the ability to detect all NTMs. The F_1 score, which is the harmonic mean of the precision and the recall, provides a balance between these two quantities (as they typically show an inverse relationship) and offers a better measure than the accuracy for an uneven class distribution.

Three approaches are considered here: the XGBoost model as above but for the binary classification; an existing heuristic model based on the *topogivity* of the elements³¹ relying only on the composition of the compounds; and a generic dimension reduction method t-SNE⁴² applied to the two most important features identified from XGBoost. All the details are available in “Methods”.

Table 2 | Comparison of the NCV accuracy, F_1 score, precision, and recall (in %) of the XGBoost, topogivity, and t-SNE approaches

	Accuracy	F_1 score	Precision	Recall	ROC AUC
XGBoost	92.4	88.5	89.5	87.5	97.5
Topogivity	87.2	79.5	85.6	74.1	90.6
t-SNE	75.7	70.8	59.0	88.3	89.1

The area under the receiver operating characteristic curve (ROC AUC) is also reported in %.

Fig. 4 | Comparison of the XGBoost, topogivity, and t-SNE approaches. Receiver operating characteristic (ROC) and precision-recall curves for distinguishing nontrivial topological materials (NTMs) from trivial insulators (TrIs) on the dataset $M \cup T$.

The results obtained on the dataset $M \cup T$ are provided in Table 2 and Fig. 4. Table 2 shows the results of the Boolean predictions with the default threshold for each algorithm.

XGBoost shows the best performance with the highest accuracy, F_1 score, precision and ROC AUC, thanks to its usage of a high-dimensional feature space to represent materials that fully describes the properties of materials. Figure 4 shows the trade-off of the scores as a function of the chosen threshold. XGBoost always has a better score. The topogivity and t-SNE approaches present an intersection point where they achieve the same scores. While their scores are lower than those of XGBoost, the topogivity and t-SNE approaches still provide reasonable results, and their advantage lies in their simplicity, making them easy to interpret.

The topogivity approach makes predictions based on a simple composition rule (see Eq. (1) in “Methods”) based on a single parameter, the elemental topogivity τ_E which approximately represents the inclination to form an NTM. Figure 5 shows a periodic table with our newly trained topogivities for 83 elements (compared to 54 available previously).

The t-SNE approach developed here focuses on two features: the maximum packing efficiency in % (MPE)⁴⁵ and the fraction of p valence electrons in % (FPV)^{46,47}. The points of the whole dataset are represented by two values representing their projections onto the t-SNE variables, as shown in Fig. 6. If the points are colored according to their type, a clear separation appears between NTMs and TrIs (in orange and blue, respectively). Taking the vertical line where t-SNE 1 is equal to zero as the splitting criterion between NTMs and TrIs, it is possible to predict 75.7% of materials correctly and to detect 88.3% of the NTMs. Furthermore, using a soft-margin linear SVM to identify the best frontier (dashed red line), the accuracy reaches 84.7%. This is still a bit lower than with the XGBoost and topogivity approaches, but it shows that even without using the target value (hence, in an unsupervised approach), the model can find the underlying relations between features and the topology of materials. The two selected features are clearly important to determine the topology of materials.

The distributions of their values in the dataset $M \cup T$ are displayed in Fig. 7. Panel (a) shows that the structures of NTMs are generally more closely packed than TrIs. This is consistent with our intuition that close-packing structures have stronger interatomic interactions, wider bands, and

higher symmetry, thus promoting the appearance of nontrivial topological phases. Panel (b) demonstrates that NTMs tend to have a lower fraction of p valence electrons. This can be rationalized as follows. Compounds with a higher fraction of p valence electrons are mainly composed of elements of the top-right part of the periodic table which are more electronegative. These tend to form ionic or strongly covalent bonds with a large trivial band gap, hence to generate TrIs. This observation aligns well with the trends in the element topogivity, as depicted in Fig. 5. Elements located in the top-right part of the periodic table display negative topogivities, indicating their inclination to form TrIs.

Discussion

In this work, a dataset of 35,608 materials with their topological properties is constructed by combining the DFT results of Materials and the Topological Materials Database, through a careful cleaning and curation process. The data from the two databases are found to be generally consistent with only 1% of the predictions which disagree. To the best of our knowledge, this is the first integration of materials from distinct data sources, a development that paves the way for more comprehensive and profound machine learning research. Using this newly created database, two research objectives were pursued.

First, machine-learning approaches were developed for categorizing materials into five distinct topological types (TrI, HSPSM, HSLSM, TI, and TCI). A thorough benchmark was performed on one of the databases to compare various machine learning approaches, obtained by combining five different models (MEGNet, Automatminer, MODNet, Random Forest, and XGBoost) with two possible procedures, namely one consisting of a series of binary-classification steps to obtain the final sorting, and the other being a direct multiclass categorization. The direct multiclass procedure relying on XGBoost was identified as the most promising approach, achieving an impressive 85.2% accuracy. A series of generalization tests were conducted that allowed for the identification of a series of differences between the two datasets (the distributions of the topological types and the chemical elements in the datasets, the presence of compounds of unknown magnetic type, as well as their coverage of the feature space). Their influence on the performance of the model was carefully analyzed.

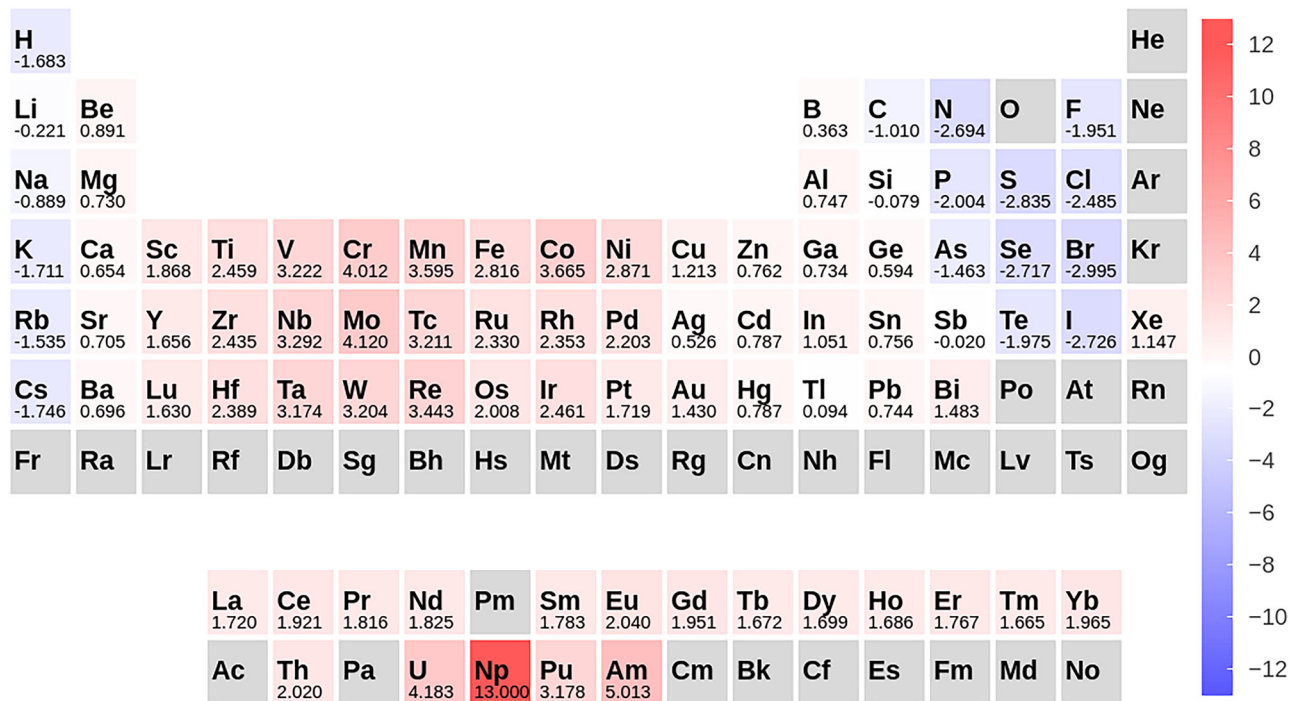


Fig. 5 | Periodic table of topogivities trained from dataset $M \cup T$. Existing topogivities are represented through numerical values with color coding, others are displayed in gray.

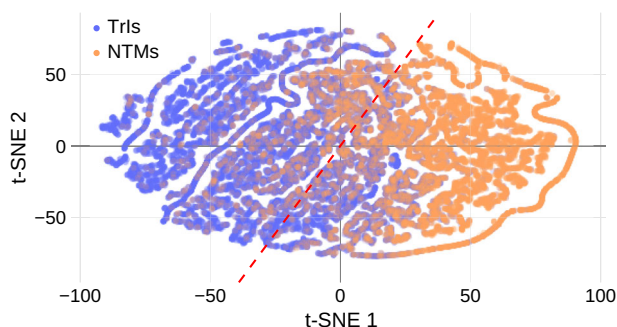


Fig. 6 | Visualization of the t-SNE results on the dataset $M \cup T$. Non-trivial materials are shown in blue, while trivial insulators are shown in orange. The red dashed line represents the decision boundary obtained using SVM.

Secondly, key factors influencing the material topology were identified by focusing on the binary classification between TrIs and NTMs. The previously developed approach relying on XGBoost performs even better on this simpler task achieving 92.4% accuracy. It was compared with two simpler methods, one relying on the use of the *topogivity* of the elements³¹, and one being an unsupervised t-SNE. The latter only focuses on the two features identified as the most important, namely MPE and FPV. It demonstrates an accuracy of 84.7%. Such performance shows that the two features are very relevant to determining the topology of materials.

Upon analyzing the distribution of these features, we found notable disparities between NTMs and TrIs. These factors are compatible with our understanding and, together with topogivity, they offer heuristic intuitions for designing topological materials. This highlights the potential to discover critical features using machine learning approaches.

Prior to our work, Claussen et al.²⁶ had used a similar machine learning approach on TMD, also relying on XGBoost. They had found that the topology is mostly determined by the chemical composition and the crystal symmetry and that it does not depend much on the particular positions of atoms in the crystal lattice. In contrast, Andrejevic et al.³⁰ had suggested that

XANES, a spectrum strongly related to the atomic type, the site, and the short-range interactions with surrounding atoms, could be used to identify NTMs and TrIs. Topogivity had also been introduced³¹ as an intuitive chemical parameter related only to the elemental composition and had been found to work rather well. Therefore, it remained elusive whether or not local chemical environment and element-related atomic features are the important characteristics influencing the topology. In our study, we included new descriptors related to the crystal structure, the composition, and the atomic sites, in addition to those employed by Claussen et al.²⁶, enlarging the space to be explored. We observed that, rather than the crystal symmetry, the most important characteristic influencing the topology is MPE of atoms in the crystal lattice space. We rationalize this by the fact that the latter actually determines the hopping parameters of electrons and as a result influences the band width and the possibility of band inversion leading to non-trivial band topology. We found that the second most important characteristic is FPV. We think that the latter has an impact on the type of bonding. Indeed, many compounds showing essentially *p* orbitals in the valence tend to be large band gap covalent insulators like diamond and silicon, or ionic insulators like NaCl and CaF₂.

Our study reveals the critical role of a comprehensive database in the ML research. The acquisition of a more extensive dataset, encompassing not only symmetry-indicator-based topological materials but also simulation results from Wilson loops, along with experimental data, holds immense importance for driving the further progress of this research.

Methods

ML models

In this study, three main approaches have been considered for the classification of topological materials. The first one is a ML classifier into the five different types (TrI, HSPSM, HSLSM, TI, and TCI). For this approach, we tested five different models combined with two possible procedures. The other two can only separate the materials into two classes, namely TrIs and NTMs (the first approach can obviously also produce this simpler classification). The second one relies on the use of a previously developed heuristic model based on the *topogivity* of the elements³¹. The last one is an unsupervised ML approach relying on t-SNE.

Fig. 7 | Distinctions of the materials in the dataset $M \cup T$. Distinction between trivial insulators (TrI) and nontrivial topological materials (NTM) based on **a** the maximum packing efficiency (%) and **b** the fraction of p valence electrons in the dataset $M \cup T$. The NTM cannot be further discriminated into HSPSM, HSLSM, TCI, and TI.

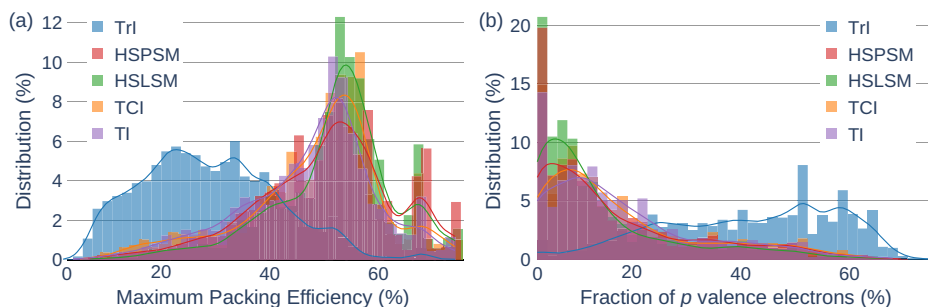


Table 3 | The set of Matminer featurizers to generate the numerical descriptors in MODNet

Composition	Structure	Site
AtomicOrbitals	BondFractions	AGNIFingerprints
AtomicPackingEfficiency	ChemicalOrdering	AverageBondAngle
BandCenter	CoulombMatrix	AverageBondLength
ElectronegativityDiff	DensityFeatures	BondOrientationalParameter
ElementFraction	EwaldEnergy	ChemEnvSiteFingerprint
ElementProperty	GlobalSymmetryFeatures	CoordinationNumber
("magpie" preset)		
IonProperty	MaximumPackingEfficiency	CrystalNNFingerprint
Miedema	RadialDistributionFunction	GaussianSymmFunc
OxidationStates	SineCoulombMatrix	GeneralizedRadialDistributionFunction
Stoichiometry	StructuralHeterogeneity	LocalPropertyDifference
TMetalFraction	XRPowderPattern	OPSiteFingerprint
ValenceOrbital		VoronoiFingerprint
YangSolidSolution		

They are used for algorithm MODNet and XGBoost.

For the ML classifier, we benchmark five different models (MEGNet, Automatminer, Random Forest, MODNet, and XGBoost) with two possible procedures. These are benchmarked on the dataset MAT to determine the best model.

MEGNet³⁴ is a graph neural network for machine learning molecules and crystals in materials science. MEGNet v1.2.9 is used in this work. For the multiclass classification (Tree 2, see below), the last layer is changed to softmax and the loss function to "categorical_crossentropy". In the MEGNet model, the crystal is represented through a CrystalGraph which is truncated using a cutoff radius of 4 Å for defining the neighbors of each atom. It is trained using 500 epochs. Given that structures containing isolated atoms cannot be handled by MEGNet, they are discarded from the training and test sets. The scores reported for MEGNet refer to the results obtained on the valid structures (i.e., without isolated atoms).

In all the other ML models, the crystal is represented using the features generated by Matminer⁴⁸. This library transforms any crystal (based on their composition and structure) into a series of numerical descriptors with a physical and chemical meaning. It relies on various featurizers adapted from scientific publications.

Automatminer (AMM)³⁵ allows for the automatic creation of complete machine learning pipelines for materials science. Here, features are automatically generated with Matminer and then reduced. Using the Tree-based Pipeline Optimization Tool (TPOT) library⁴⁹, an AutoML stage is used to prototype and validate various internal ML pipelines. A customized 'express' preset pipeline is performed on the training set (note that the "EwaldEnergy" is excluded from the AutoFeaturizer due to technical problems).

Random Forest (RF)³⁹, which is an ensemble learning method, constructs multiple decision trees. For the classification task, the final output is the result of majority voting. Here, we first use Matminer to extract Magpie_ElementProperty⁴⁶, SineCoulombMatrix⁵⁰, DensityFeatures and GlobalSymmetryFeatures from the input structures (the missing features are filled with the average of the known data). The hyperparameters are determined using a 5×3 NCV on the training data, where the stratified 3-fold inner cross-validation uses a grid search for determining the optimal values (`{'max_features': ['auto', 'sqrt', 'log2'], 'criterion': ['gini', 'entropy']}`) relying on the 'balanced_accuracy' score.

MODNet^{36,41} is a framework based on feature selection and a feed-forward neural network. The selection uses a relevance-redundancy score defined from the mutual information between the features and the target and between pairs of features. The framework is well suited for limited datasets. Here, we first use a predefined set of featurizers (DeBreuck2020Featurizer with accelerated oxidation state parameters) to generate features from the structures. Table 3 provides a complete list of all these Matminer featurizers. The missing features are filled with their default value (most are zero). The feature selection is performed on the training data only. The model hyperparameters are determined using a genetic algorithm.

XGBoost⁴⁰ is an ensemble of boosted trees. Here, the features generated by the MODNet approach are used as the input of the model. The hyperparameters are determined using a 5×5 NCV on the training data, where the stratified 5-fold inner cross-validation uses a grid search for determining the optimal values (see Table 4) relying on the F_1 score.

Two different classification procedures are used, as shown in Fig. 8. The first approach (Tree 1) uses four one-vs.-all binary-classification steps to obtain the final classification. The second approach (Tree 2) uses a direct

multi-class classification for the five different types, using either majority voting for tree-based methods, or softmax activation for neural network based methods. The accuracy for both approaches are reported in the figure.

A consistent NCV testing procedure is used to evaluate the performance on the dataset *MAT* for the ten different combinations of ML models and procedures. For all the ML models, an identical stratified 5-fold outer loop was used to determine the generalization error (test error).

The internal validation depends on the algorithm, as explained above.

In terms of the final classification, the results of the two procedures are very close with a ranking that depends on the ML model. The best result is achieved for the direct multiclass procedure relying on XGBoost which achieves an accuracy of 85.2%.

The topogivity τ_E of an element E has been proposed as a measure of its tendency to form topological materials³¹. The τ_E values for 54 elements were originally obtained by using support vector machine (SVM)³³ model trained on a subset of 9026 compounds from the database created by Tang et al.⁸, with approximately one half classified as trivial and the other half as topological. The ML model is based on a heuristic chemical rule, which maps

each material M to a number $g(M)$ through the function:

$$g(M) = \sum_E f_E(M) \tau_E \quad (1)$$

where the summation runs over all elements E in the chemical formula of material M , with $f_E(M)$ denoting the fraction of element E within material M . A material M is classified based on $g(M)$, as trivial (TrI) if negative and topological (NTM) if positive.

Here, we also build a new topogivity model trained on a subset from dataset $M \cup T$, which excludes the elements occurring less than 25 times. We construct a soft-margin linear SVM using the scikit-learn library (specifically, the `sklearn.svm.SVC` class). The hyperparameter C of the model is determined through a grid search among 15 values evenly spaced on a log scale ranging from 10^4 to 10^6 relying on the F_1 score and adopting a 5-fold validation procedure.

The optimal value ($C = 3.73 \times 10^4$) is then used to train the final model on the whole dataset. This new topogivity model provides the τ_E value for 83 elements and covers 35,522 materials out of the 35,608 of the dataset $M \cup T$. These values are reported on the periodic table shown in Fig. 5.

For comparison, the previous model³¹ gives the τ_E value for 54 elements and covers only 18,637 compounds in $M \cup T$. A quantitative comparison of the two models can thus only be performed on these 18,367 compounds. From the different scores reported in Table 5, it can be said that their performance is essentially similar.

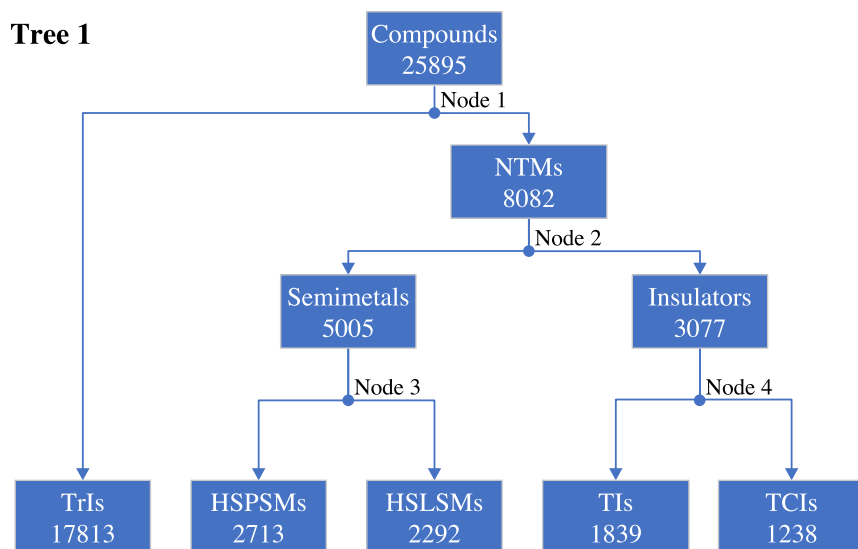
To compare the differences element by element, we use a form of relative difference defined by:

$$\Delta\tau_E = \frac{2 \operatorname{sign}(\tau_E^p \tau_E^n) |\tau_E^p - \tau_E^n|}{2 + |\tau_E^p| + |\tau_E^n|}, \quad (2)$$

Table 4 | Hyperparameter grid searched for the XGBoost model

Parameter	Value	Parameter	Value
Learning rate η	0.23	L^2 regularization λ	1.33
Maximal tree depth	[9, 10, 11]	Minimil child weight	[0.1, 0.3, 0.5]
Column subsampling by tree	[0.75, 0.78]	Column subsampling by node	[0.75, 0.78]

Tree 1



Tree 2

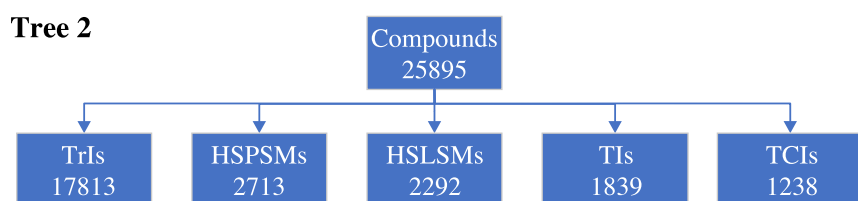
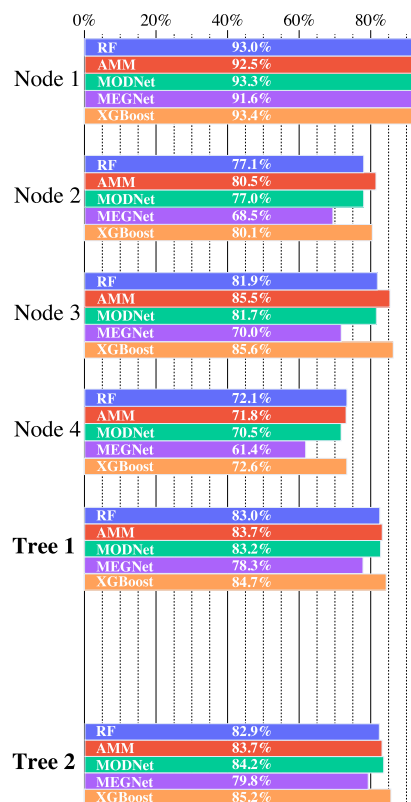


Fig. 8 | Illustration of the two different procedures used to categorize the materials according to their topological types. The Tree 1 relies on four binary-classification steps, while Tree 2 is based on a direct multiclass approach. The accuracy

achieved with the five ML models (RF, AMM, MODNet, MEGNet, and XGBoost) are reported both at the final stage of both procedures, as well as at each step in Tree 1.



where τ_E^p and τ_E^n are the topogivities of the previous and new models, respectively. The values of $\Delta\tau_E$ are indicated on the periodic table shown in Fig. 9. They show that the two sets of results are essentially consistent with each other, except for three elements (Li, Si, and Sb), which exhibit different signs (as indicated by the blue color in the figure).

The final evaluation of the topogivity approach is performed using a 5×5 NCV. The results are presented in Table 2 and Fig. 4. It leads to an accuracy of 87.2%, a F_1 score of 79.5%, a precision of 85.6%, and a recall of 74.1%.

The last approach originates from the idea to identify the most important features to distinguish topological materials, which we then combine with the unsupervised algorithm t-SNE. All the Matminer features are first ranked in descending order of gain importance for training the XGBoost classification model between the five types for the datasets $M \cup T$, $M \cap T$ and $T \cup M$. The intersection of the top 15 most important features consists of three features: the maximum packing efficiency (MPE) in %, the fraction of p valence electron (FPV) in %, and the formation enthalpy (ΔH) in eV/atom as obtained from the semi-empirical Miedema model⁵¹.

Based on this finding, the distribution of the compounds in $M \cup T$ according to the values of MPE, FPV, and ΔH is analyzed to understand what helps the models to distinguish between TrIs from NTMs. The corresponding plots, in which the five different materials types have been separated, are reported in Fig. 7 for MPE and FPV (already commented in the “Results” section) and in Fig. 10a for ΔH .

For the latter, it turns out that there is a concentration of TrIs around zero. The reason for this is, however, not physical but technical. In fact, the semi-empirical Miedema model⁵¹ does not apply to all materials. For those

compounds where it does not work, Matminer does not produce a value for the feature and, as commonly done in ML approaches, we replace these missing values with a zero. As it turns out, as shown in Fig. 10c, the share of such missing values is much larger for TrIs (more than 87%) than NTMs, the ML model takes advantage of this flaw to identify them.

If the compounds without ΔH value are removed, 12,518 entries are left in $M \cup T$, rather evenly sorted into 3622 TrIs, 2843 HSPSMs, 2734 HSLSMs, 1339 TCIs, and 1980 TIs (see Fig. 10d). The corresponding distribution of the ΔH values is reported in Fig. 10b. It still shows a significant difference between TrIs (which mostly present positive values) and NTMs (which are mainly negative). The underlying reasons are still not completely clear to us. Given that value of the feature is missing for many materials, it can, however, not be used as a discriminator. Therefore, it is discarded from the subsequent analysis.

The dataset $M \cup T$ can now be simply visualized in 2D by representing each material by its values for MPE and FPV, as shown in Fig. 11. The plot reveals a rather clear distinction between TrIs and NTMs. And, if we apply t-SNE on the whole dataset, the distinction is even clearer. Finally, based on the t-SNE variables, a dividing line can be drawn using a soft-margin linear SVM in which C is equal to 1.0.

Heterogeneity metric

To quantify the degree of diversity within a dataset or between two datasets, we adopted a heterogeneity metric defined as the k -nearest-neighbor (k -NN) distance in the space of the most important Matminer features. The most important features are defined based on their induced gain during the training of the XGBoost model. A total of 47 Matminer features (44 presenting continuous values and 3 discrete ones, respectively) was gathered by taking the union of the 20 most important ones for training on the sub-datasets $M \cup T$, $M \cap T$, and $T \cup M$. These originate from the five featurizers reported in Table 6.

All the continuous feature values are rescaled to range between 0 to 1 using the MinMaxScaler from scikit-learn on the dataset $M \cup T$. Then, the distance $d(x, y)$ between any two points x and y is defined as:

$$d(x, y) = \sqrt{\sum_{i=1}^{n_c} (x_i - y_i)^2 + \sum_{j=n_c+1}^{n_c+n_d} (1/2(x_j == y_j))^2}, \quad (3)$$

Table 5 | NCV accuracy, F_1 score, precision and recall (in %) of the two different topogivity models using the 18,637 compounds of the $M \cup T$ dataset which only include the 54 elements for which τ_E is provided in ref. 31

	Accuracy	F_1 score	Precision	Recall
This work	90.2	76.8	83.5	71.2
Ref. 31	89.6	77.3	77.1	77.5

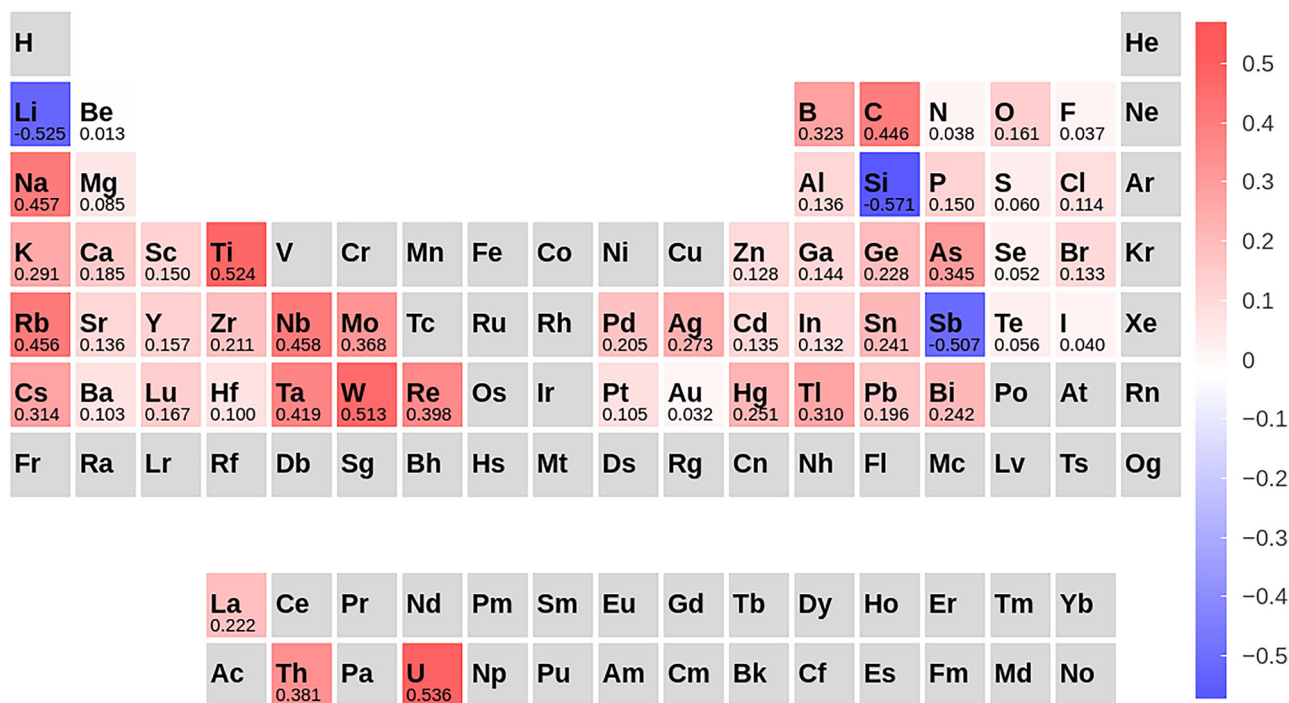


Fig. 9 | Comparison of the topogivities of ref. 31 with those obtained here for 54 elements. The color code refers to $\Delta\tau_E$, a measure of their relative difference, as defined by Eq. (2). When the signs of the two topogivities are different, the value of $\Delta\tau_E$ is negative, and thus the element is highlighted in blue.

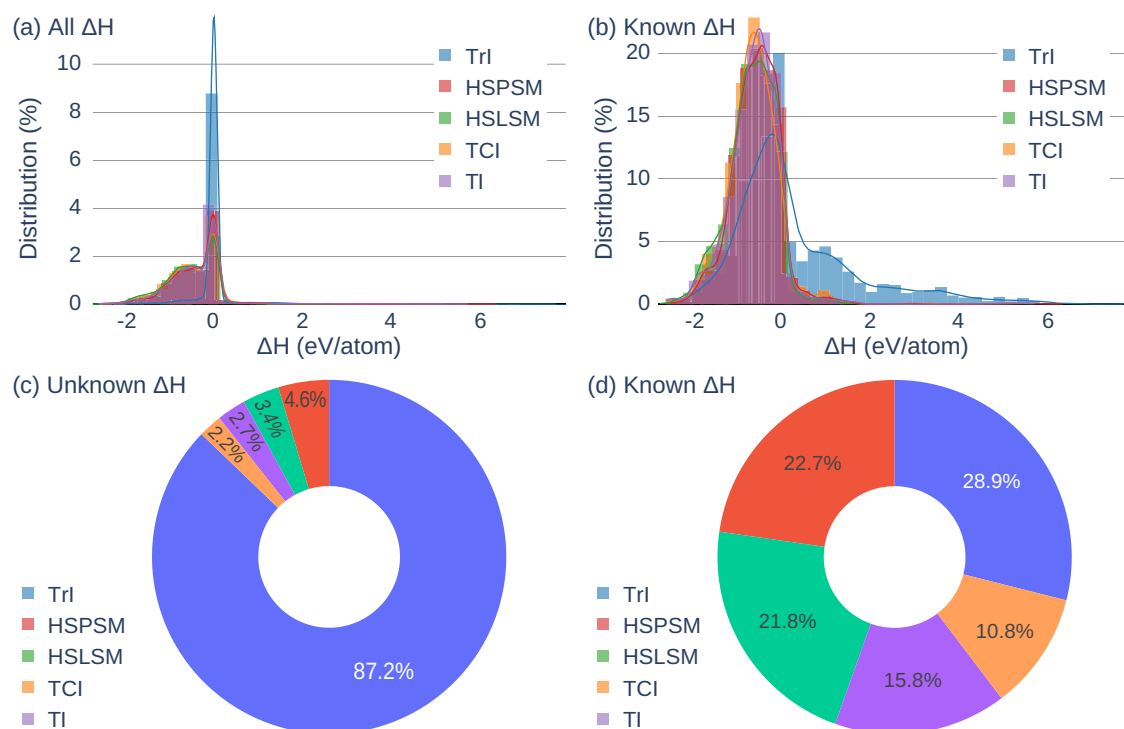


Fig. 10 | Distribution of ΔH values in the dataset $M \cup T$ according to the five topological types. **a** Shows the distribution of all the values, while **b** presents the distribution only for those compounds for which the value is known through the

Miedema model⁵¹. **c**, **d** Show the pie-charts with distribution of the compounds among the five different topological types for those compounds with unknown and known values, respectively.

Table 6 | Matminer featurizers generating the features selected for defining the heterogeneity metric between datasets

Matminer featurizer	Description
• Miedema	Formation enthalpies of intermetallic compounds ⁵¹ .
• ElementProperty ("magpie" preset)	Weighted elemental statistics ⁴⁶ .
• OxidationStates	Statistics about the oxidation states for each specie. Features are concentration-weighted statistics of the oxidation states.
• AtomicOrbitals	Estimation of the highest occupied molecular orbital (HOMO) and lowest unoccupied molecular orbital (LUMO) energies based on the composition and the atomic orbital energies ⁵⁶ .
• GlobalSymmetryFeatures	Symmetry information.

A brief description and the source of the data is also provided.

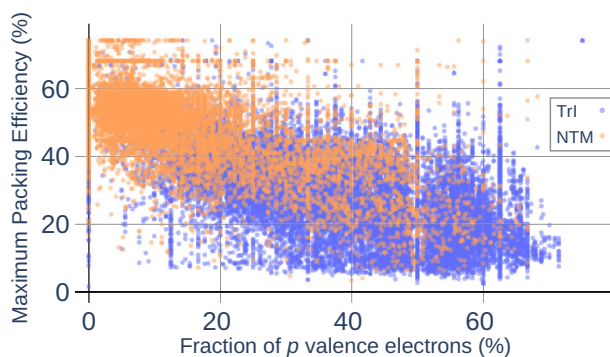


Fig. 11 | Visualization of the dataset $M \cup T$ according to the two most important features: the maximum packing efficiency and the fraction of p valence electrons. This representation leads to a clear separation between TrIs and NTMs even though there is some overlap between the two types.

where $n_c = 44$ and $n_d = 3$ are the numbers of features presenting continuous and discrete values, respectively. The distance $D(x, A)$ between a point x and a dataset A is defined as the mean of the distances from point x to the 5 NN-points (q_j with $j = 1, \dots, 5$) in dataset A :

$$D(x, A) = \frac{1}{5} \sum_{j=1}^5 d(x, q_j). \quad (4)$$

Finally, the heterogeneity metric $H(A, B)$ between two datasets A and B (which can be the same dataset A) is defined based on the average distance between all the points of A (p_i with $i = 1, \dots, N_A$) and the dataset B :

$$H(A, B) = \frac{1}{N_A} \sum_{i=1}^{N_A} D(p_i, B). \quad (5)$$

Note that with this definition $H(A, B)$ is an asymmetric measure (it follows from the fact that training and testing sets are not interchangeable).

Data availability

The data can be downloaded from MatElab (<https://in.iphy.ac.cn/eln/link.html#119/P6g5>) and the Materials Cloud Archive⁵² <https://doi.org/10.24435/materialscloud:zk-gc>. It can be viewed interactively (e.g., the plots corresponding to Figs. 7, 11) through the chemscope visualization tool⁵³. It can be queried using the OPTMADE API^{54,55} using the following link: <https://optmade.materialscloud.org/archive/zk-gc/v1/info>. A Jupyter notebook is also available. All this information is also available here: <https://topoclass.modl-uclouvain.org>.

Received: 25 February 2025; Accepted: 2 June 2025;

Published online: 13 June 2025

References

- Kitaev, A. Periodic table for topological insulators and superconductors. In *AIP Conference Proceedings*, Vol. 1134, 22–30 (American Institute of Physics, 2009).
- Hasan, M. Z. & Kane, C. L. Colloquium: topological insulators. *Rev. Mod. Phys.* **82**, 3045 (2010).
- Qi, X.-L. & Zhang, S.-C. Topological insulators and superconductors. *Rev. Mod. Phys.* **83**, 1057 (2011).
- Po, H. C., Vishwanath, A. & Watanabe, H. Symmetry-based indicators of band topology in the 230 space groups. *Nat. Commun.* **8**, 50 (2017).
- Bradlyn, B. et al. Topological quantum chemistry. *Nature* **547**, 298–305 (2017).
- Zhang, T. et al. Catalogue of topological electronic materials. *Nature* **566**, 475–479 (2019).
- Vergniory, M. G. et al. A complete catalogue of high-quality topological materials. *Nature* **566**, 480–485 (2019).
- Tang, F., Po, H. C., Vishwanath, A. & Wan, X. Comprehensive search for topological materials using symmetry indicators. *Nature* **566**, 486–489 (2019).
- Kruthoff, J., De Boer, J., Van Wezel, J., Kane, C. L. & Slager, R.-J. Topological classification of crystalline insulators through band structure combinatorics. *Phys. Rev. X* **7**, 041069 (2017).
- Song, Z., Zhang, T., Fang, Z. & Fang, C. Quantitative mappings between symmetry and topology in solids. *Nat. Commun.* **9**, 3530 (2018).
- Khalaf, E., Po, H. C., Vishwanath, A. & Watanabe, H. Symmetry indicators and anomalous surface states of topological crystalline insulators. *Phys. Rev. X* **8**, 031070 (2018).
- Hohenberg, P. & Kohn, W. Inhomogeneous electron gas. *Phys. Rev.* **136**, B864–B871 (1964).
- Kohn, W. & Sham, L. J. Self-consistent equations including exchange and correlation effects. *Phys. Rev.* **140**, A1133–A1138 (1965).
- Hellenbrandt, M. The inorganic crystal structure database (ICSD)-present and future. *Crystallogr. Rev.* **10**, 17–22 (2004).
- Watanabe, H., Po, H. C. & Vishwanath, A. Structure and topology of band structures in the 1651 magnetic space groups. *Sci. Adv.* **4**, aat8685 (2018).
- Elcoro, L. et al. Magnetic topological quantum chemistry. *Nat. Commun.* **12**, 5965 (2021).
- Peng, B., Jiang, Y., Fang, Z., Weng, H. & Fang, C. Topological classification and diagnosis in magnetically ordered electronic materials. *Phys. Rev. B* **105**, 235138 (2022).
- Samuel, A. L. Some studies in machine learning using the game of checkers. *IBM J. Res. Dev.* **3**, 210–229 (1959).
- Zhang, Y. & Kim, E.-A. Quantum loop topography for machine learning. *Phys. Rev. Lett.* **118**, 216401 (2017).
- Zhang, P., Shen, H. & Zhai, H. Machine learning topological invariants with neural networks. *Phys. Rev. Lett.* **120**, 066401 (2018).
- Zhang, Y., Ginsparg, P. & Kim, E.-A. Interpreting machine learning of topological quantum phase transitions. *Phys. Rev. Res.* **2**, 023283 (2020).
- Scheurer, M. S. & Slager, R.-J. Unsupervised machine learning and band topology. *Phys. Rev. Lett.* **124**, 226401 (2020).
- Donoho, D. L. Compressed sensing. *IEEE Trans. Inf. Theory* **52**, 1289–1306 (2006).
- Friedman, J. H. Greedy function approximation: a gradient boosting machine. *Ann. Stat.* **29**, 1189–1232 (2001).
- Acosta, C. M. et al. Analysis of topological transitions in two-dimensional materials by compressed sensing. Preprint at *arXiv* <https://doi.org/10.48550/arXiv.1805.10950> (2018).
- Claussen, N., Bernevig, B. A. & Regnault, N. Detection of topological materials with machine learning. *Phys. Rev. B* **101**, 245117 (2020).
- Cao, G. et al. Artificial intelligence for high-throughput discovery of topological insulators: the example of alloyed tetradymites. *Phys. Rev. Mater.* **4**, 034204 (2020).
- Liu, J., Cao, G., Zhou, Z. & Liu, H. Screening potential topological insulators in half-heusler compounds via compressed-sensing. *J. Phys. Condens. Matter* **33**, 325501 (2021).
- Schleder, G. R., Focassio, B. & Fazzio, A. Machine learning for materials discovery: two-dimensional topological insulators. *Appl. Phys. Rev.* **8**, 031409 (2021).
- Andrejevic, N. et al. Machine-learning spectral indicators of topology. *Adv. Mater.* **34**, 2204113 (2022).
- Ma, A. et al. Topogivity: a machine-learned chemical rule for discovering topological materials. *Nano Lett.* **23**, 772–778 (2023).
- Vergniory, M. G. et al. All topological bands of all nonmagnetic stoichiometric materials. *Science* **376**, eabg9094 (2022).
- Cortes, C. & Vapnik, V. Support-vector networks. *Mach. Learn.* **20**, 273–297 (1995).
- Chen, C., Ye, W., Zuo, Y., Zheng, C. & Ong, S. P. Graph networks as a universal machine learning framework for molecules and crystals. *Chem. Mater.* **31**, 3564–3572 (2019).
- Dunn, A., Wang, Q., Ganose, A., Dopp, D. & Jain, A. Benchmarking materials property prediction methods: the Matbench test set and automatminer reference algorithm. *npj Comput. Mater.* **6**, 138 (2020).
- De Breuck, P.-P., Hautier, G. & Rignanese, G.-M. Materials property prediction for limited datasets enabled by feature selection and joint learning with MODNet. *npj Comput. Mater.* **7**, 83 (2021).
- Xie, T. & Grossman, J. C. Crystal graph convolutional neural networks for an accurate and interpretable prediction of material properties. *Phys. Rev. Lett.* **120**, 145301 (2018).
- Stone, M. Cross-validated choice and assessment of statistical predictions. *J. R. Stat. Soc. Ser. B Methodol.* **36**, 111–133 (1974).
- Ho, T. K. Random decision forests. In *Proceedings of 3rd International Conference on Document Analysis and Recognition*, Vol. 1, 278–282 (IEEE, 1995).
- Chen, T. & Guestrin, C. XGBoost: a scalable tree boosting system. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining. KDD '16*, 785–794 (ACM, 2016).
- De Breuck, P.-P., Evans, M. L. & Rignanese, G.-M. Robust model benchmarking and bias-imbalance in data-driven materials science: a case study on modnet. *J. Phys. Condens. Matter* **33**, 404002 (2021).
- Van der Maaten, L. & Hinton, G. Visualizing data using t-SNE. *J. Mach. Learn. Res.* **9**, 2579–2605 (2008).
- Jain, A. et al. Commentary: The Materials Project: a materials genome approach to accelerating materials innovation. *APL Mater.* **1**, 011002 (2013).
- Ong, S. P. et al. Python materials genomics (pymatgen): a robust, open-source python library for materials analysis. *Comput. Mater. Sci.* **68**, 314–319 (2013).
- Ward, L. et al. Including crystal structure attributes in machine learning models of formation energies via Voronoi tessellations. *Phys. Rev. B* **96**, 024104 (2017).
- Ward, L., Agrawal, A., Choudhary, A. & Wolverton, C. A general-purpose machine learning framework for predicting properties of inorganic materials. *npj Comput. Mater.* **2**, 16028 (2016).

47. Deml, A. M., O'Hayre, R., Wolverton, C. & Stevanović, V. Predicting density functional theory total energies and enthalpies of formation of metal-nonmetal compounds by linear regression. *Phys. Rev. B* **93**, 085142 (2016).
48. Ward, L. et al. Matminer: an open source toolkit for materials data mining. *Comput. Mater. Sci.* **152**, 60–69 (2018).
49. Olson, R. S. et al. *Applications of Evolutionary Computation: 19th European Conference, EvoApplications 2016, Porto, Portugal, Proceedings, Part I, Chap. Automating Biomedical Data Science Through Tree-Based Pipeline Optimization*, 123–137 (Springer International Publishing, 2016).
50. Faber, F., Lindmaa, A., von Lilienfeld, O. A. & Armiento, R. Crystal structure representations for machine learning models of formation energies. *Int. J. Quantum Chem.* **115**, 1094–1101 (2015).
51. Zhang, R., Zhang, S., He, Z., Jing, J. & Sheng, S. Miedema calculator: a thermodynamic platform for predicting formation enthalpies of alloys within framework of miedema's theory. *Comput. Phys. Commun.* **209**, 58–69 (2016).
52. Talirz, L. et al. Materials cloud, a platform for open computational science. *Sci. Data* **7**, 299 (2020).
53. Fraux, G., Cersonsky, R. & Ceriotti, M. Chemscope: interactive structure-property explorer for materials and molecules. *J. Open Source Softw.* **5**, 2117 (2020).
54. Andersen, C. W. et al. OPTIMADE, an API for exchanging materials data. *Sci. Data* **8**, 217 (2021).
55. Evans, M. L. et al. Developments and applications of the optimade API for materials discovery, design, and data exchange. *Digit. Discov.* **3**, 1509–1533 (2024).
56. Kotochigova, S., Levine, Z. H., Shirley, E. L., Stiles, M. D. & Clark, C. W. Local-density-functional calculations of the energy of atoms. *Phys. Rev. A* **55**, 191–199 (1997).

Acknowledgements

We thank N. Regnault for providing the data of TMD, which makes the data curation of the Materiae and TMD possible. Y.H. and H.W. acknowledge the funding from the National Key Research and Development Program of China (Grant No. 2022YFA1403800) and the National Natural Science Foundation of China (Grant No. 12188101). Y.H. was supported by China Scholarship Council (Grant No. 201904910878). H.W. is also supported by the New Cornerstone Science Foundation through the XPLOER PRIZE. Computational resources have been provided by the supercomputing facilities of the Université catholique de Louvain (CISM/UCL) and the Consortium des Équipements de Calcul Intensif en Fédération Wallonie Bruxelles (CÉCI) funded by the Fond de la Recherche Scientifique de Belgique (F.R.S.-FNRS) under convention 2.5020.11 and by the Walloon Region.

Author contributions

Y.H. wrote the main manuscript text. She performed the data curation, integration, and worked on the machine learning model development and generalization tests. P.D. contributed to model optimization and hyperparameter tuning. H.W. provided critical insights into the topological properties of materials and supervised the theoretical aspects of the study. Matteo Giantomassi contributed to the analysis of the features. Gian-Marco Rignanese supervised the project and guided the research direction. All authors discussed the results, reviewed the manuscript, and approved the final version.

Competing interests

The authors declare no competing interests.

Additional information

Supplementary information The online version contains supplementary material available at <https://doi.org/10.1038/s41524-025-01687-2>.

Correspondence and requests for materials should be addressed to Hongming Weng or Gian-Marco Rignanese.

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

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

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

© The Author(s) 2025