

Title	Automated SAT problem feature extraction using convolutional autoencoders
Authors	Dalla, Marco;Visentin, Andrea;O'Sullivan, Barry
Publication date	2021-11-01
Original Citation	Dalla, M., Visentin, A. and O'Sullivan, B. (2021) 'Automated SAT Problem Feature Extraction using Convolutional Autoencoders', ICTAI 2021: IEEE International Conference on Tools with Artificial Intelligence Virtual event, 01-03 November, pp. 232-239. doi: 10.1109/ICTAI52525.2021.00039
Type of publication	Conference item
Link to publisher's version	https://ieeexplore.ieee.org/document/9643220 - 10.1109/ICTAI52525.2021.00039
Rights	For the purpose of Open Access, the authors have applied a CC BY public copyright licence to this Author Accepted Manuscript; Copyright published version: © 2021 IEEE; - https://creativecommons.org/licenses/by/4.0/
Download date	2026-09-25 02:25:40
Item downloaded from	https://hdl.handle.net/10468/12238

Automated SAT Problem Feature Extraction using Convolutional Autoencoders

Marco Dalla

*SFI Centre for Research Training
in Artificial Intelligence
School of Computer Science & IT
University College Cork, Ireland
m.dalla@cs.ucc.ie*

Andrea Visentin

*Confirm Centre for Smart Manufacturing
School of Computer Science & IT
University College Cork, Ireland
andrea.visentin@ucc.ie*

Barry O’Sullivan

*Insight Centre for Data Analytics
School of Computer Science & IT
University College Cork, Ireland
b.osullivan@cs.ucc.ie*

Abstract—The Boolean Satisfiability Problem (SAT) is the first known NP-complete problem and has a very broad literature focusing on it. It has been applied successfully to various real-world problems, such as scheduling, planning and cryptography. SAT problem feature extraction plays an essential role in this field. SAT solvers are complex, fine-tuned systems that exploit problem structure. The ability to represent/encode a large SAT problem using a compact set of features has broad practical use in instance classification, algorithm portfolios and solver configuration. The performance of these techniques relies on the ability of feature extraction to convey helpful information. Researchers often craft these features “by hand” to capture particular structures of the problem. Instead, in this paper, we extract features using semi-supervised deep learning. We train a Convolutional Autoencoder (AE) to compress the SAT problem in a limited latent space and reconstruct it minimizing the reconstruction error. The latent space projection should preserve much of the structural features of the problem. We compare our approach to a set of features commonly used for algorithm selection. Firstly, we train classifiers on the projection to predict if the problems are satisfiable or not. If the compression conveys valuable information, a classifier should be able to take correct decisions. In the second experiment, we check if the classifiers can identify the original problem that was encoded as SAT. The empirical analysis shows that the autoencoder is able to represent problem features in a limited latent space efficiently, as well as convey more information than current feature extraction methods.

1. Introduction

A SAT instance consists of a Boolean formula that involves a set of Boolean variables connected by the logical operators “and”, “or” and “not”. The Propositional Boolean Satisfiability problem consists of finding an assignment to the variables such that the formula evaluates to true (satisfiable instance) or proves that such an assignment does not exist (unsatisfiable instance). SAT solvers are software designed to take such instances and find a solution or prove that one does not exist. Many studies have been dedicated to developing new solving techniques and improving existing

ones. These solvers compete annually in the SAT competition [1] to test the state of the art in this discipline on hard problem instances. These solvers use different solving techniques, such as Conflict-Driven Clause Learning (CDCL), Look-Ahead SAT, and Stochastic Local Search. In these competitions it is clear that no solver outperforms all others on all problems; depending on the instance type, their rank varies consistently. For example, CDCL solvers are especially efficient in solving industrial SAT instances. The reason behind this specialization is not yet clear. Due to this, portfolio algorithms often perform best. A portfolio in this setting is a collection of individual solvers that can be selectively deployed depending on the type of instance that has to be tackled. Portfolios deploy a multi-class classifier on the instance features to select the best solver. This task is called the Algorithm Selection Problem [2].

Another approach to tackle a SAT instance is to predict the satisfiability as a classification problem using machine learning. The first attempt in this field [3] uses features extracted by a portfolio algorithm. Recently, the pervasiveness of deep learning extended to this field as well. The most relevant attempt is NeuroSAT [4], a graph deep learning network that iteratively tries to classify a graphical representation of the SAT instance. This approach can be seen as a deep learning heuristic solver.

A fundamental component of these approaches is the extraction of meaningful features from a SAT instance, e.g. [5]. This task aims to represent a problem that involves a high number of variables and clauses into a handful of values preserving the relevant information on its structure and properties. The features are generally statistical information of the instance, e.g. the number of variables, clauses or their ratio, or computed on different representations of the problem, such as the graph encoding. The classic process of extracting features mirrors the human view of the problem structure. One of the reasons why it is hard to understand the behaviour of the solvers on different types of problems is that we are trying to see them from a human perspective. Machine-automated feature extraction can reduce the human bias on the problem description.

In this work, we present a semi-supervised deep learning automated approach to extract features from SAT instances.

We train a convolutional autoencoder (AE) to compress and reconstruct an instance minimizing the reconstruction error. An AE is a deep network that takes the input and learns to copy it, minimizing the differences between the original and the network’s output. The particularity is that one of the internal layers has a limited number of neurons, acting as an information bottleneck. The AE learns to compress the relevant information to fit through this bottleneck and reconstruct them correctly. It can be decomposed into an encoder and a decoder; the encoder projects the input into the latent space represented by the bottleneck, while the decoder takes a point of the latent space and tries to reconstruct the original instance. This can be seen as a lossy compression technique. During the training phase, the AE learns to a low dimensional representation of the input; if this representation preserves the information, a classifier should be able to take decisions on these features.

2. Literature review

We survey the relevant literature. We focus on feature extraction techniques for SAT instances and SAT/UNSAT classification. Finally, we briefly present some successful example of AE used in feature extraction.

In their seminal paper, Mitchell et al. [6] show a strong correlation between the hardness of a SAT instance and its ratio between the number of variables and clauses. Building on that work, researchers have attempted to extract from SAT problems a number of features that measure, with increasing precision, their empirical hardness and, therefore, how likely a solver will be able to evaluate its satisfiability within a meaningful runtime. Nudelman et al. [7] introduced 84 features that can be calculated from SAT instances. These characteristics can be split into nine categories: problem size, variable-clause-graph, variable graph, clause graph, balance, proximity to Horn formulae, LP-based, DPLL-probing, and local search probing. One of the first works that uses these features to build an effective empirical hardness model was SATzilla-07 [8]. The model is used then to construct a per-instance algorithm portfolio that automatically selects among different SAT solvers to minimize the runtime of the solver on the instance. Malitzky et al. [9] proposed a new systematic approach for calculating features that do not require hand-crafted structural features to guide the selection of the algorithm. They define these as latent features stemming from matrix decomposition. By feeding those to a linear model, they reach performance comparable to modern algorithm portfolios. Moreover, the information gathered from the latent features can help researchers extract only the necessary features that would guide the algorithm selection. Their projection differs from the work presented herein since they start from the features pre-computed by other approaches, while our AE takes the SAT instance as input. Finally, two recent papers by Ansótegui et al. [5] and Matos Alfonso et al. [10] propose two new sets of features that further improve the choice for a specific algorithm to solve specific SAT instances. In the first, they concentrate specifically on structural characteristics of industrial samples

(i.e. scale-free structure, community structure, and self-similar structure). They then train a classifier that assesses the effectiveness of these new sets of features compared to previously introduced SAT features sets. The second presents three new features based on the structural information of the formula and the consideration AND-gates as well as exactly-one constraints. Likewise, they use these features to train a random forest algorithm that guides the selection of the appropriate SAT solver configuration.

As shown previously, Boolean Satisfiability can be approached and aided through machine learning and deep learning methods. One of the first examples of using these techniques is demonstrated in a 2008 paper by Devlin et al. [3]. They use multiple machine learning classifiers to directly predict SAT instance satisfiability by training them on the same features used by the SATzilla algorithm portfolio. Another important paper was published by Grozea et al. in 2014 [11]. They employ machine learning models to correctly predict branches during the search for a solution for 3-SAT instances. They find that in over 90% of the test cases using a machine learning heuristic-enhanced SAT solver reduces by at most 1/3 the number of branchings that the algorithm performed during search. In 2017 Wu [12] suggests an improvement to the runtime of the MiniSAT solver. By utilising a logistic regression model to predict the satisfiability of Boolean formulae with a certain number of fixed variables, he applies the learned model to secure a preferable initial value for each Boolean variable using a Monte-Carlo approach. The prediction algorithm manages to guess 78% of the backbones correctly. This approach reduces the number conflicts in the search for a solution, therefore improving the runtime. However, the total runtime needs to consider the preprocessing phase in which the logistic regression model is trained. Consequently, it is not able to outperform overall the vanilla Minisat. In a 2016 paper [13], Loreggia et al. investigate a new approach using deep learning to select the preferred algorithm for solving specific SAT and CSP instances. After encoding the ASCII representation of the problem into a black and white image, they trained a deep convolutional neural network to classify between different solvers. This is based on the (strong) assumption that only a specific solver can assess the satisfiability with one particular runtime. They observe that on a test set, the network is able to predict the correct algorithm with over 90% accuracy. In a 2017 paper [14], Bunz et al. implement, for the first time, a deep learning algorithm to predict the satisfiability of instances. After providing a graphical representation of conjunctive normal forms stemming from a hand-generated random 3-SAT dataset, they feed the representation directly to a graph neural network. They manage to achieve over 65% accuracy on the testing set without any previous feature extraction. Finally, a 2018 paper [?], [4] by Selsam et al. introduces Neurosat, a message parsing neural network that learns to provide a solution to SAT instances by only being trained to predict their satisfiability. While this approach cannot yet outperform state of the art SAT solvers, it is able to scale efficiently with the size of the problem and handle instances

with completely different structures from the ones it is trained on. Feature extraction is one of the many applications of AEs. Dong et al. [15] provide a structured description of the technique and a comprehensive review of their literature. The non-linear projection computed by an AE is widely used to compress data efficiently into a low-dimensionality latent space that can be considered a set of features of the original data. For example, in a 2011 paper [16], Masci et al. use stacked convolutional autoencoders to extract biologically plausible features of samples from two famous datasets, MNIST and CIFAR-10. Then these features are used to train a convolutional neural network that performs an efficient classification. Liang et. al. [17] survey the literature in text feature extraction techniques based on deep learning. This technique is currently employed in many other fields, like semiconductor manufacturing [18], sensory robotics [19], radar data analysis [20], etc. The surveyed literature shows the research interest for the task presented in this paper and the potential that the AE shows in similar tasks in different domains.

3. Encoding

Conjunctive normal form (CNF) is the instance format used by the vast majority of the solvers. A CNF is composed of a conjunction of *clauses*. Each clause is a disjunction of *literals*. A literal is the occurrence of a variable negated or not. All SAT problems can be transformed into CNF formulas in linear time with only a linear increase in the formula size [21]. The DIMACS CNF format is the standard, however, it is not well suited for an AE. Different approaches have been introduced, such as transforming it into an image [13] or using a graph representation [14]. The network has to output a multi-class classification. One-hot encoding is the general approach for such problems.

Our approach transforms a CNF into an equivalent binary array. Given a CNF with C clauses and V variables, we encode each clause i as an array of binary variables $\mathbf{x}^i$. For each variable j , $x_{2j}^i = 1$ if j appears in clause i , and $x_{2j+1}^i = 1$ if j appears negated in that clause, otherwise they take value 0. Thereby each CNF can be encoded in $2CV$ binary variables. We fixed a maximum number of clauses and variables encoding all instances according to those sizes to make input size constant.

3.1. Symmetry breaking

A drawback of CNF formulas is that they are subjected to different invariances. These symmetries make the network learning process harder since the network needs to learn that different inputs might encode an equivalent state of the problem. Graph representation removes some of these symmetries [22]. Invariances are particularly detrimental in our case due to the loss computation. A reconstructed instance that represents the same problem but with variables and clauses that are permuted has a really high reconstruction loss. In our case, we considered an ordering that reduces the invariances of our encoding.

Many operations can change the CNF without affecting its satisfiability. We tackled them in the following way:

- *Variable negation*, for example, negating all the occurrences of a given variable. For each variable, the normal literal has to appear at least as much as the negated one.
- *Permuting variables*, for example swapping all the occurrences of variable $j = 1$ with variable $j = 2$. We ordered them by the number of the literal appearances. The first variable is the one that appears the most in the formula. Ties are broken considering the variable with the minimum amount of negated literals.
- *Permuting literals in a clause*, changing the order of the literals in a clause does not affect the solution. Our encoding is not affected by this symmetry.
- *Permuting clauses*, for example, swapping the first clause with the second one. We order the clauses by their cardinality (number of literals appearing in the clause). Ties are broken considering the literal with the lower index.

While this approach does not resolve some invariances completely, it allows a network to learn patterns. For example, the variable that appears the most appears negated in a clause with a variable that rarely appears.

3.2. Dataset

The main reason for generating the dataset is the necessity for a sizeable number of instances with diverse structures that are well balanced between SAT and UNSAT. Deep learning solution quality strongly depends on the quantity and how representative the training dataset. Moreover, the use of deep learning algorithms and the binary encoding presented in the previous section makes handling large instances quite complicated in terms of memory usage; therefore, a reduced number of variables and clauses was required. To the best of our knowledge, no publicly available dataset is satisfying these constraints, so we opted for generating one. Creating a generator for problem instances that have the structure and computational properties that are more similar to real-world instances has been an ongoing challenge [23]. The dataset used in the empirical section of this paper is generated using CNFgen [24]. This tool produces propositional formulas in the CNF DIMACS format that can be used as a benchmark for SAT solvers. It features several formula families (e.g. pigeonhole principle, ordering principle, k-coloring, etc.) as well as several formula transformations and the possibility of producing formulas directly from graph structures.

Two datasets were assembled. The first one, which we will later refer to as the training dataset, consists of 3000 instances uniformly distributed in eight types, mainly stemming from graph structures. Of these 3000 instances, 1562 are proven unsatisfiable and 1438 satisfiable beforehand using the Glucose SAT solver [25]. This set is used to train the AE in order to find the optimal representation in the

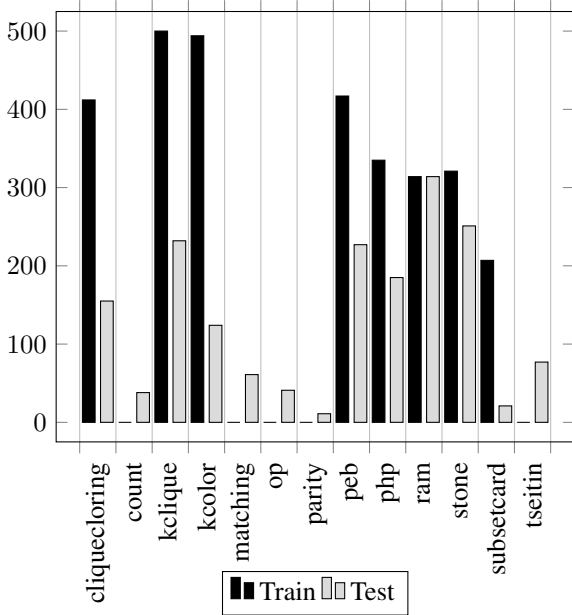


Figure 1: Distribution of the different SAT instances for the training dataset and testing dataset

latent space. The learned embedding is then used to train a machine learning classifier to predict both the satisfiability and the instance type.

A second dataset, for testing, is generated to assess how well the autoencoder performs on previously unseen instances. This set consists of 1737 instances divided into 12 classes. Out of those 12 classes, 8 are of the same type used for the training dataset, while the remaining 5 are new types of instances. As previously, we determine the satisfiability of the instances using the Glucose solver, resulting in 1035 unsatisfiable samples and 702 satisfiable samples. In the empirical experiments, the training set is used as a whole, while the test set is considered both as a whole and divided between previously seen and new instance types. This separation aims to analyse the ability of the AE to generalise to unseen data structures. The distribution of both datasets is shown in Figure 1.

4. Method

The deep learning algorithm chosen for the feature extraction is a Convolutional Autoencoder (AE) [26]. It is a type of feed-forward neural network architecture that leverages the technique of representation learning. Specifically, the network is composed of two branches that function as an encoder-decoder pair. At the conjunction, a bottleneck is present to obtain a compressed knowledge representation of the input. The bottleneck compression can be referred to as a latent space of lower dimensionality into which the input is projected. This technique is based on the assumption that the input data has a structure that can be efficiently learned and used in its reduced form to reproduce the input. Indeed, the task of an AE is to generate an output that is a

reconstruction of the original input. The **reconstruction loss** measures how well the original input has been reproduced in the output from the latent space. In this case, the chosen reconstruction loss was the *binary-crossentropy* [27]; this loss considers every binary variable of the encoding as a binary classification task. This choice is possible due to the binary encoding of the input presented in the previous section. Before feeding the input into the network, it undergoes the encoding and symmetry breaking process.

Using convolutional layers offers many advantages over a dense AE. These layers are widely used in image and signal processing. They can recognize patterns even when their position is moved around the input. We can use them to learn literals and clause patterns. Convolutional layers can process bigger inputs compared to fully connected layers. Finally, it can focus the scope of a literal to the clause it is part of. In previous works using convolutional neural networks, such as [13], a window includes literals from consecutive clauses, learning patterns on those.

The encoder part of the network takes the input and projects it into the latent space. It starts with a convolutional layer with a single filter that encodes a literal, with windows of size 2 that do not overlap. This layer converts each literal to a single value. We prefer this approach over using a three-valued encoding for each variable in each clause (such as 0 for non-present, 1 for present, and 2 for present and negated) because this encoding would imply different errors in the case of the case misclassification. The second convolutional layer has non-overlapping windows with a size equal to the maximum number of variables. Each window covers all the literals involved in a clause without overlapping to literals in other clauses. The advantage is that it can learn patterns at clause level, e.g. clauses that involve the most used variables all negated. Finally, we deploy a set of convolutional layers. Using a series of convolutional layers helps to learn higher-level patterns, and it generally outperforms both in training times and performances a single layer with more filters. The encoder ends with a short series of fully connected layers, the last one with a number of neurons equal to the latent space size.

These layers, with the exception of the final one, are followed by dropout [28] and batch-normalization [29] layers. These make the algorithm less sensitive to overfitting, reduce the internal covariant shift and allow a higher learning rate without the possibility of gradient explosion or vanish.

The decoder mirrors the encoder structure with transposed convolutions. We select a latent space of size $\delta = \{8, 16, 24, 32, 40\}$. The optimizer Adam [30] is used to speed up the back-propagation. The initial learning rate is set to 10^{-4} with a decay of 10^{-5} . Finally, the Rectified Linear Unit [31] is chosen as the activation function. Figure 2 shows the entire structure of the AE implemented and used in the experiments.

Before the choosing of the optimal parameters and the AE structure, a large hyperparametrisation was conducted. We trained all the networks for multiple days on three NVIDIA Quadro RTX 8000 graphic cards.

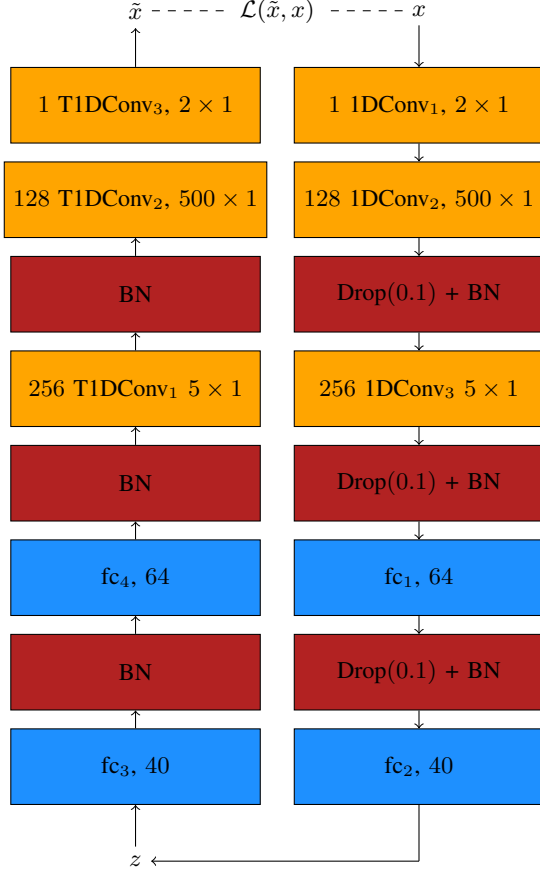


Figure 2: Convolutional Autoencoder consisting of encoder (right) and decoder (left). The encoder features 3 convolutional layers, followed by 2 fully connected layers. It also features dropout and batch normalization layers after the second and third convolutional layers and the first dense layer. The decoder mirrors the encoder structure, with transpose convolutions instead. For the convolutional layers, the number of filters and window sizes are specified, for the dense layer the number of neurons. The reconstruction loss $\mathcal{L}(\tilde{x}, x)$ quantifies the quality of the reconstruction and is minimized during training.

5. Experimental results

The empirical evaluation aims to compare the features extracted using the approach presented in this paper to the widely used SatZilla features [8]. SATZilla is a portfolio SAT solver that for each instance selects the solver based on a set of 127 features that includes different aspects of the problem: size, graph, balance and timing features. Of this set of features we ignored the timing/probing ones that are solver dependent and not specific of the CNF. The goal of a compression/feature extraction solution is to preserve/extract relevant information. To assess the importance of the information extracted a common approach is to train a classifier on the features; if a classifier achieves a high accuracy it

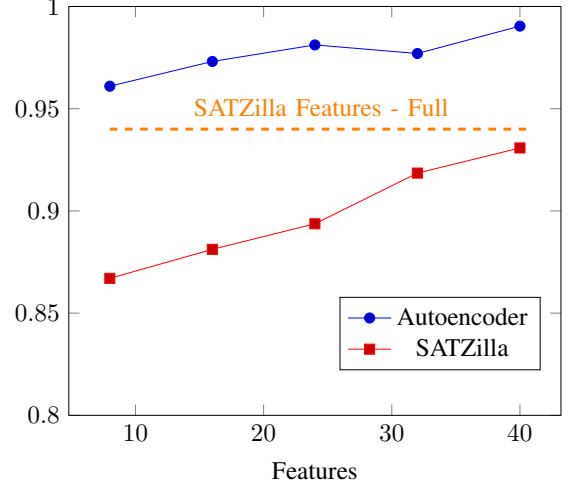


Figure 3: SAT/UNSAT classification accuracy (training set).

means that the features extract enough information to take a correct decision. Firstly, we check if the features can be used for the binary classification of the satisfiability of the problem. In a second experiment, we build a classifier that is able to recognize the type of graph problem encoded. We use the dataset presented in Section 3.2. We train the AE using an 80/20 validation split on the train dataset; we then use it to extract the features of both train and test set. Since SATZilla does not have a training phase its deployment is equivalent on the train and test set. We presents the results of XGBoost [32] as classifier. Then, the average accuracy is computed over 10 shuffles of 10 cross validation. We also analyse the performance on the training set. The AE is optimised for the instance compression, not for the classification task; comparing its performances with the test set we can understand how well the features generalise to unseen structures.

5.1. SAT/UNSAT classification

In this experiment, we try to predict the satisfiability of a problem based on the extracted features. This task is widely studied in the literature, see Section 2. We project all the instances of the dataset into increasing size latent spaces, $\delta = \{8, 16, 24, 32, 40\}$. The goal is to test the ability of the approach to compress the instance structure. We train AEs for the different δ values on the training set. For the SATZilla features, we compress them using principal component analysis (PCA) to reduce their cardinality.

Figure 3 and Figure 4 show the results of our experiment on the train and the test set, respectively. The orange dotted line is the accuracy achieved by using the full SATZilla features set. The AE outperforms the best SATZilla in the training set even when compressing the SAT instance to just eight features, while in the test set that contains also different types of problems is comparable. When projecting to 40 dimensions the gap between the two increases.

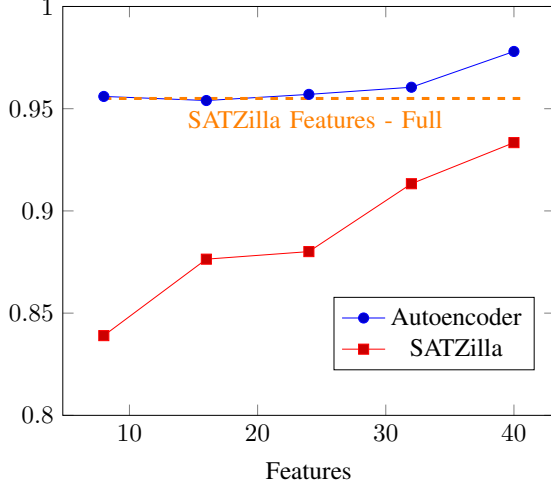


Figure 4: SAT/UNSAT classification accuracy (test set).

Figure 5 shows the confusion matrices for the best and worse feature extraction settings of the two approaches on the test set. It is clear that the PCA compression of the SATZilla features reduces the ability of the classifier to recognize the UNSAT instances. For the AE, when increasing the dimension of the latent space the accuracy improves similarly in both classes. We have similar results in the training set, where the best AE classify all the SAT instances correctly and makes mistakes in just a few UNSAT ones. To evaluate if our approach can generalise to new types of instances we divided the test set in instance types that are present in the training set as well, and the ones that are not. As Figure 6 shows, the AE outperforms SATZilla in both cases. The new instances are likely harder to classify since both of them have worse performances, but this could be due to the smaller size of the dataset as well.

To understand if these differences are statistically significant, we study the performances of the best solutions of the two approaches using the Bayesian classifier comparison presented in [33]. This analysis allows us to compute the posterior probability of a classifier being better than the other and the probability of them being equivalent from a practical point of view (1% accuracy of difference). Figure 7 shows the results. The AE approach has the 96.73% probability to be better than SATZilla in SAT/UNSAT classification on the test set, and the 3.27% of being practically equivalent. For the sake of brevity, we omitted the plot computed on the training set where the AE has a 99.99% probability of being practically better than the competitor.

The algorithm presented in this work computes features that preserve the instance satisfiability information in a better way compared to the existing approaches. Increasing the compression level marginally affect the performances of the AE, but it has a consistent impact on the SATZilla features. The ability to extend these results to unseen instances with different structure (originated from different problems) shows the generalization abilities of SATZilla features.

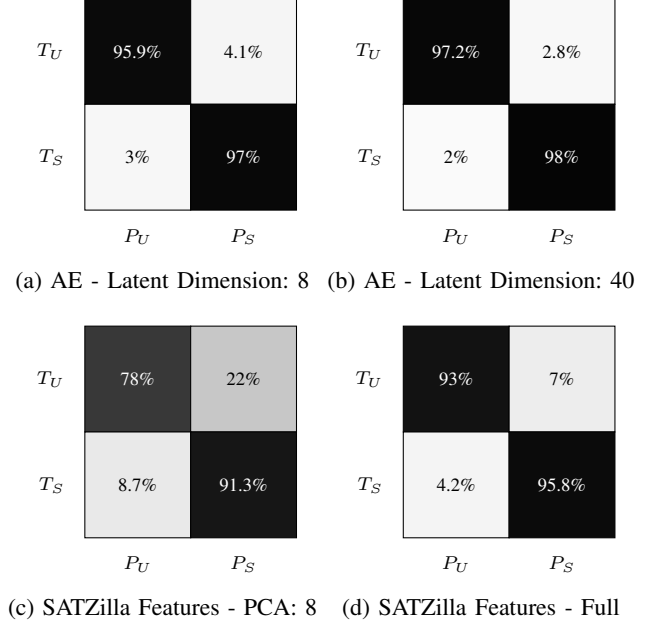


Figure 5: Confusion matrices for the test set using different dimension of the latent space for the AE and SATZilla. T_U and T_S are respectively the true UNSAT and the true SAT label, while P_U and P_S are the predicted values.

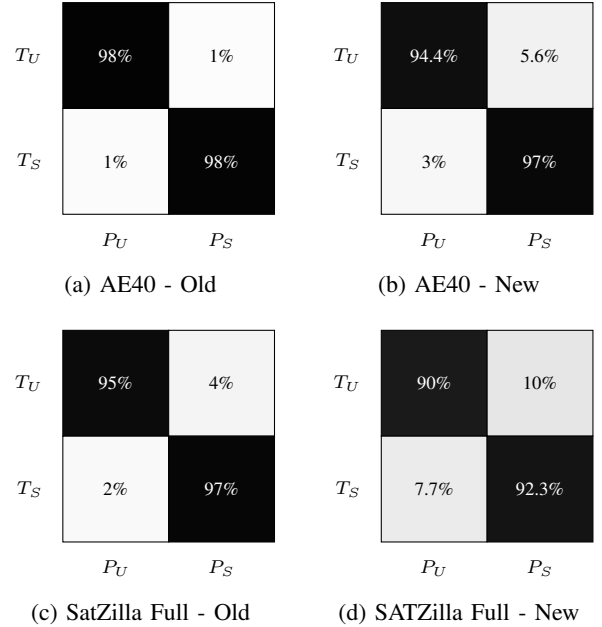


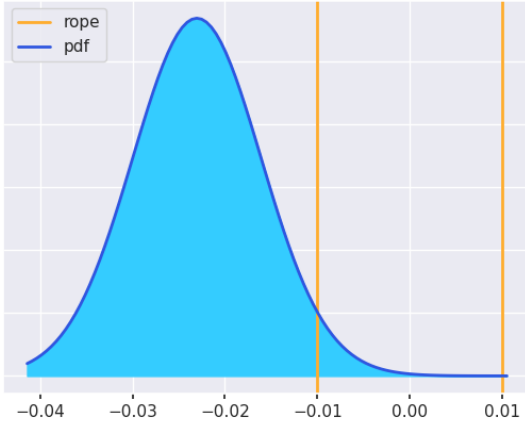
Figure 6: Confusion matrices for the test set split into two groups: the first with similar types of instances of the ones used to train the AE, the second with new types of instances.

5.2. Instance classification

One of the most efficient approaches to solve SAT instances is using portfolio solvers. These solvers use a collection of normal SAT solvers and select a specific one

TABLE 1: instance classification accuracy

Instance Type	Training set		Test set - Old		Test set - New		Test set - Full	
	Autoencoder	SATZilla	Autoencoder	SATZilla	Autoencoder	SATZilla	Autoencoder	SATZilla
cliquecoloring	73	95	58	76	-	-	58	91
count	-	-	-	-	50	67	50	64
kclique	86	97	73	80	-	-	67	89
kcolor	95	100	67	86	-	-	79	93
matching	-	-	-	-	86	75	96	87
op	-	-	-	-	100	88	100	80
parity	-	-	-	-	98	91	98	87
peb	93	100	90	100	-	-	96	100
php	80	95	97	100	-	-	89	100
ram	100	100	100	100	-	-	100	100
stone	97	97	100	97	-	-	88	89
subsetcard	50	38	65	63	-	-	93	85
tseitin	-	-	-	-	75	100	75	97
Average	85.9	93.4	86.1	94.3	82.6	91.8	84.5	93.3

Figure 7: Posterior distribution of the comparison between AE 40 features and SATZilla full feature set. *rope* is the region of practical equivalence.

depending on the instance. SATZilla features are designed for this purpose, while AE features aim to represent the whole instance in a limited projected space. In this experiment, we tried to identify the original problem class encoded as a CNF instance using the classes presented in [24]. We used the same features extracted in the previous experiment.

Table 1 shows the classification accuracy divided by class for the AE with 40 latent dimensions and the complete set of SATZilla features. The classes not present on the set are left empty. We analyse the test set complete, and divided between old and new instance types. SATZilla outperforms the AE in this task. We assume this is because some instances have similar structures but differ for some statistical features. For example, the AE struggles to distinguish the graph-based problems (cliquecoloring, kclique,

kcolor). Likely, these problems have similar structures, and their differences are not well represented in the projected space. For some instances, the AE has better performances; for the training set in subset-cardinality, and for the test set in matching, op, parity. SATZilla features are better for identifying the type of instance. We believe that these problems are better represented by the compression approach we use.

6. Conclusions

In this paper, we presented a new automated approach to extract features from a SAT instance. In our approach, the features are learned by an unsupervised neural network and not crafted by a human. To do so, we used a convolutional autoencoder designed to compress and reconstruct the SAT instance minimizing the reconstruction error. To minimize this error, the compressed space (latent projection) has to preserve the information on the original instance. The deep neural network exploits a CNF encoding in binary variables and a symmetries breaking approach we introduced.

The empirical study shows that our approach can preserve the instance information to allow a classifier to predict the satisfiability of the problem and the type of instance. Compared to the state of the art, our approach can convey more information in a limited feature space. The difference between the two approaches is statistically significant and relevant from a practitioner point of view. SATZilla features outperform the presented approach on instance type classification. However, the results show that the information conveyed by the AE can help to identify types of instances in which the statistical features struggles.

Our approach has performance relevant for practitioners. However, its applicability is still limited by the overhead it introduces. The encoding increases the space required by a CNF and, for large instances, the time required to train the network. This problem is common in the vast majority

of deep learning approaches in SAT solving, such as the famous NeuroSAT. They allow us to obtain a different point of view of this well-studied problem.

Our work opens a variety of research questions. A deep analysis of the latent space can correlate its dimensions to the artificial features, this might allow us to discover blind spots in the statistical features currently in use. We plan to improve the results further by introducing reconstruction error losses tailor-made for SAT instances. Another aspect we want to cover is creating a set that comprises both human-designed and deep learning computed features to encode different aspects of the SAT instances. Finally we intend to investigate the possibility of using deep learning generative models such as variational autoencoders.

Acknowledgments. This publication has emanated from research conducted with the financial support of Science Foundation Ireland under Grant number 18/CRT/6223. For the purpose of Open Access, the author has applied a CC BY public copyright licence to any Author Accepted Manuscript version arising from this submission.

References

- [1] T. Balyo, N. Froleyks, M. J. Heule, M. Iser, M. Järvisalo, and M. Suda, “Proceedings of sat competition 2020: Solver and benchmark descriptions,” 2020.
- [2] J. R. Rice, “The algorithm selection problem,” *Adv. Comput.*, vol. 15, pp. 65–118, 1976. [Online]. Available: [https://doi.org/10.1016/S0065-2458\(08\)60520-3](https://doi.org/10.1016/S0065-2458(08)60520-3)
- [3] D. Devlin and B. O’Sullivan, “Satisfiability as a classification problem,” *Proc. of the 19th Irish Conf. on Artificial Intelligence and Cognitive Science*, 2008.
- [4] D. Selsam, M. Lamm, B. Bünz, P. Liang, L. de Moura, and D. L. Dill, “Learning a SAT solver from single-bit supervision,” in *Proceedings of ICLR 2019*. OpenReview.net, 2019. [Online]. Available: https://openreview.net/forum?id=HJMC_iA5tm
- [5] C. Ansótegui, M. L. Bonet, J. Giráldez-Cru, and J. Levy, “Structure features for sat instances classification,” *Journal of Applied Logic*, vol. 23, pp. 27–39, 2017.
- [6] D. G. Mitchell, B. Selman, and H. J. Levesque, “Hard and easy distributions of SAT problems,” in *Proc. of AAAI*, 1992, pp. 459–465.
- [7] E. Nudelman, K. Leyton-Brown, H. H. Hoos, A. Devkar, and Y. Shoham, “Understanding random SAT: beyond the clauses-to-variables ratio,” in *Proceedings of CP*, ser. Lecture Notes in Computer Science, M. Wallace, Ed., vol. 3258. Springer, 2004, pp. 438–452. [Online]. Available: https://doi.org/10.1007/978-3-540-30201-8_33
- [8] L. Xu, F. Hutter, H. H. Hoos, and K. Leyton-Brown, “Satzilla: Portfolio-based algorithm selection for SAT,” *JAIR*, vol. 32, pp. 565–606, 2008.
- [9] Y. Malitsky and B. O’Sullivan, “Latent features for algorithm selection,” in *Proceedings of SOCS*, 2014.
- [10] E. M. Alfonso and N. Manthey, “New CNF features and formula classification,” in *Proceedings of POS Workshop*, D. L. Berre, Ed., vol. 27, 2014, pp. 57–71.
- [11] C. Grozea and M. Popescu, “Can machine learning learn a decision oracle for np problems? a test on sat,” *Fundamenta Informaticae*, vol. 131, pp. 441–450, 2014.
- [12] H. Wu, “Improving sat-solving with machine learning,” in *Proceedings of ACM SIGCSE*, 2017, pp. 787–788.
- [13] A. Loreggia, Y. Malitsky, H. Samulowitz, and V. A. Saraswat, “Deep learning for algorithm portfolios,” in *Proceedings of AAAI*, 2016, pp. 1280–1286.
- [14] B. Bünz and M. Lamm, “Graph neural networks and boolean satisfiability,” *arXiv*, pp. 1–9, 2017.
- [15] G. Dong, G. Liao, H. Liu, and G. Kuang, “A review of the autoencoder and its variants: A comparative perspective from target recognition in synthetic-aperture radar images,” *IEEE Geoscience and Remote Sensing Magazine*, vol. 6, no. 3, pp. 44–68, 2018.
- [16] J. Masci, U. Meier, D. C. Cireşan, and J. Schmidhuber, “Stacked convolutional auto-encoders for hierarchical feature extraction,” in *Proceedings of ICANN*, vol. 6791, 2011, pp. 52–59.
- [17] H. Liang, X. Sun, Y. Sun, and Y. Gao, “Text feature extraction based on deep learning: a review,” *EURASIP journal on wireless communications and networking*, vol. 2017, no. 1, pp. 1–12, 2017.
- [18] M. Maggipinto, C. Masiero, A. Beghi, and G. A. Susto, “A Convolutional Autoencoder Approach for Feature Extraction in Virtual Metrology,” *Procedia Manufacturing*, vol. 17, pp. 126–133, 2018.
- [19] M. Polic, I. Krajacic, N. Lepora, and M. Orsag, “Convolutional Autoencoder for Feature Extraction in Tactile Sensing,” *IEEE Robotics and Automation Letters*, vol. 4, no. 4, pp. 3671–3678, 2019.
- [20] H. Lee, J. Kim, B. Kim, and S. Kim, “Convolutional autoencoder based feature extraction in radar data analysis,” in *Proceedings of Joint SCIS and ISIS*, 2018, pp. 81–84.
- [21] G. S. Tseitin, “On the complexity of derivation in propositional calculus,” in *Automation of reasoning*. Springer, 1983, pp. 466–483.
- [22] D. Selsam and N. Bjørner, “Guiding high-performance sat solvers with unsat-core predictions,” vol. 11628 LNCS, pp. 336–353, 2019.
- [23] H. Kautz and B. Selman, “Ten challenges redux: Recent progress in propositional reasoning and search,” in *Proceedings of CP*, 2003, pp. 1–18.
- [24] M. Lauria, J. Elffers, J. Nordström, and M. Vinyals, “Cnfgcn: A generator of crafted benchmarks,” in *Proceedings of SAT*, 2017, pp. 464–473.
- [25] G. Audemard and L. Simon, “On the glucose sat solver,” *Int. J. Artif. Intell. Tools*, vol. 27, pp. 1 840 001:1–1 840 001:25, 2018.
- [26] M. A. Kramer, “Nonlinear principal component analysis using autoassociative neural networks,” *AIChE Journal*, vol. 37, no. 2, pp. 233–243, 1991.
- [27] U. Ruby and V. Yendapalli, “Binary cross entropy with deep learning technique for image classification,” *International Journal of Advanced Trends in Computer Science and Engineering*, vol. 9, 10 2020.
- [28] N. Srivastava, G. Hinton, A. Krizhevsky, I. Sutskever, and R. Salakhutdinov, “Dropout: A simple way to prevent neural networks from overfitting,” *JMLR*, vol. 15, pp. 1929–1958, 06 2014.
- [29] S. Ioffe and C. Szegedy, “Batch normalization: Accelerating deep network training by reducing internal covariate shift,” in *Proceedings of ICML*, 2015, pp. 448–456.
- [30] D. P. Kingma and J. Ba, “Adam: A method for stochastic optimization,” in *Proceedings of ICLR*, Y. Bengio and Y. LeCun, Eds., 2015. [Online]. Available: <http://arxiv.org/abs/1412.6980>
- [31] A. F. Agarap, “Deep learning using rectified linear units (relu),” *ArXiv*, vol. abs/1803.08375, 2018.
- [32] T. Chen and C. Guestrin, “Xgboost: A scalable tree boosting system,” in *Proceedings of the 22nd acm sigkdd international conference on knowledge discovery and data mining*, 2016, pp. 785–794.
- [33] A. Benavoli, G. Corani, J. Demsar, and M. Zaffalon, “Time for a change: a tutorial for comparing multiple classifiers through bayesian analysis,” *J. Mach. Learn. Res.*, vol. 18, pp. 77:1–77:36, 2017.