

Topological Analysis of Contradictions in Text

Xiangcheng Wu

xwu20@uncc.edu

University of North Carolina at

Charlotte

Charlotte, NC, USA

Xi Niu

xniu2@uncc.edu

University of North Carolina at

Charlotte

Charlotte, NC, USA

Ruhani Rahman

rrahman3@uncc.edu

University of North Carolina at

Charlotte

Charlotte, NC, USA

ABSTRACT

Automatically finding contradictions from text is a fundamental yet under-studied problem in natural language understanding and information retrieval. Recently, topology, a branch of mathematics concerned with the properties of geometric shapes, has been shown useful to understand semantics of text. This study presents a topological approach to enhancing deep learning models in detecting contradictions in text. In addition, in order to better understand contradictions, we propose a classification with six types of contradictions. Following that, the topologically enhanced models are evaluated with different contradictions types, as well as different text genres. Overall we have demonstrated the usefulness of topological features in finding contradictions, especially the more latent and more complex contradictions in text.

CCS CONCEPTS

• **Information systems** → *Information retrieval*; • **Computing methodologies** → *Natural language processing*.

KEYWORDS

topological data analysis; deep learning; contradiction; text representation

ACM Reference Format:

Xiangcheng Wu, Xi Niu, and Ruhani Rahman. 2022. Topological Analysis of Contradictions in Text. In *Proceedings of the 45th International ACM SIGIR Conference on Research and Development in Information Retrieval (SIGIR '22)*, July 11–15, 2022, Madrid, Spain. ACM, New York, NY, USA, 6 pages. <https://doi.org/10.1145/3477495.3531881>

1 INTRODUCTION

Text data play an essential role in our lives. Since we communicate using natural languages, we produce and consume a large amount of text data every day. The explosive growth of text data makes it impossible for people to digest relevant text data in a timely manner. Typically text data contains two types of information: facts and opinions. With the flushing of fake news and misinformation, different or even conflicting versions of "facts" confuse people constantly. As to opinions, people by nature have different opinions and disagreements. Therefore, we believe finding contradictions

in text is a fundamental problem in text understanding, for both facts and opinions, and its further applications such as fake news identification, computational surprise [1, 14–16], and argument retrieval [7, 20, 23, 24], etc. Automatically identifying contradictions is a potential solution to help sense-making and decision-making process in information seeking. However, the problem of contradiction is an under-studied problem in natural language processing and information retrieval [12].

Topology is a classic branch of mathematics that concerns with the properties of shape that are preserved under continuous deformations, such as stretching, twisting, crumpling, and bending, but not tearing or gluing. Topology has recently been shown useful to understand semantics of text [26]. Although in recent years deep learning techniques have achieved huge success in understanding natural language due to their strong capacity of representing text, such high dimensional representations are usually not isotropic: they occupy a narrow cone in the vector space and are not uniformly distributed with respect to direction [21]. Meanwhile, high-dimensional data usually contains plenty of noises and therefore it is essential to ensure the feature representations are not "washed out" by these noises. Informally speaking, **Topological Data Analysis (TDA)** can be viewed as a process of data compression and representation [13]. However, TDA focuses more on the representation of data as a whole than deep learning approaches. This new perspective from topology may help us uncover a deeper and more holistic nature of data. In addition, these underlying shapes basically remain the same even if the data is slightly deformed. This means that topological features can reflect the essential characteristics of data to a certain extent and reduce the impacts of noises.

In this study, we present a topological approach to enhancing the deep learning models in detecting contradictions in text. In addition, we propose a classification with six types of contradictions from a perspective of language phenomena. Under the classification system, we manually labeled the types of contradictions to make a ground truth testing dataset and further evaluated our proposed models' performances for different types of contradictions.

The main contributions of this paper are two-fold: 1) we designed an approach to incorporating the topological representation of text into the deep learning representations. With the enhanced representation of text, we obtained improved model performance in contradiction detection; and 2) we proposed a classification with six types of contradictions from the language perspective, and manually labeled a dataset as the ground truth testing data.

2 RELATED WORK

In this section we will introduce preliminaries of **Topological Data Analysis (TDA)**. Then we will review a few studies that used TDA in text understanding research.

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than ACM must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

SIGIR '22, July 11–15, 2022, Madrid, Spain

© 2022 Association for Computing Machinery.

ACM ISBN 978-1-4503-8732-3/22/07...\$15.00

<https://doi.org/10.1145/3477495.3531881>

2.1 Preliminaries of TDA

Topological features are the “holes” in an entity. Because the “holes” cannot be created by continuous deformations, they show how entities are consistent or different with others in topology. The formation of “holes” relies on simplicial complexes, which is a set of simplexes in different dimensions [8]. An i -simplex should be a figure with a minimum number of $i - 1$ -simplexes as faces. For example, A 0-simplex, 1-simplex, 2-simplex and 3-simplex refers to a point (0-dimension), a line segment (1-dimension), a triangle (2-dimension), and a tetrahedron (3-dimension) respectively. Thus, the “holes” are defined as the empty space enclosed by the simplexes in the corresponding dimensions within an entity.

There are many different approaches to simplicial complex constructions. TDA utilizes the Vietoris-Rips (VR) filtration approach [19], which is easy for a point in a high-dimensional space to be tracked and defined in matrices.

In deep learning, a piece of text is represented in a high-dimensional space, which can be considered as an entity. Specifically, the piece of text could be represented as a series of points (each point representing a word), called a point cloud in a high-dimensional space. However, in the point cloud, there is no “hole”. We need a technique to connect these points to form “holes”. This technique in TDA is called Persistent Homology (PH) [2, 9].

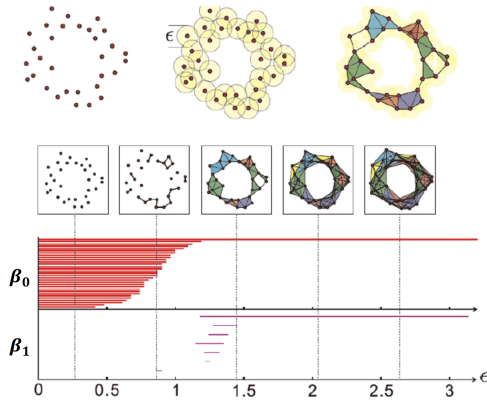


Figure 1: Persistence Diagram [11]. Once we start expanding the data points into balls of increased radii ϵ , planar figures emerge. The bars in β_0 and β_1 capture relevant features of this process, and therefore information about the original data. This is how persistence homology (the main aspect of TDA) works. This figure is cited from [11].

Figure 1, cited from the study of [11], shows the PH of text. Initially a unit of text (sentence, paragraph, document, etc) is represented as a cloud of points (each point represents a word), and later its subsequent approximation by balls of increased radius ϵ . The overlaps produce a change in shape which can be measured using the number of i -dimensional “holes”, represented as β_i [9]. β_0 ’s definition is a bit different from other β_i s, as it refers to the number of connected components of the entity. As in the β_0 and β_1 lines in Figure 1, the number of β_0 lines intersecting the vertical bar at ϵ represents the number of connected components of when the points are extended with balls of that radius. Therefore as ϵ

increases, the number of components decreases. The β_1 lines show the birth and death of the one dimensional “holes” at given values of ϵ . In this process the exact positions of the data points are ignored, but the shape the cloud is preserved. That is, two clouds of similar shapes but different positions will have similar persistence diagrams. Collectively, β_0 and β_1 (and higher β_i s, not discussed here) compress and represent information about the shape of the point cloud. This diagram only shows planar structures, but persistence patterns work in higher dimensions as well, in principle allowing machines to “see” shapes in dimensions higher than 3, a task difficult for humans.

2.2 Previous Studies on TDA for Understanding Text

A decade ago, TDA was introduced in text mining based on the observation that data points may have implicit shapes (e.g., [26]). A natural question to ask is whether texts have shapes that can be measured using tools of topology. Zhu in 2013 [26] was the first to investigate this question. Zhu used a collection of nursery rhymes to illustrate how topology can be used to find certain patterns of repetition. More recently, Doshi et al. [6] applied Zhu’s method in a larger setting showing its classification superiority on the task of assigning movie genres to user generated plot summaries for the IMDB dataset. In another study, Savle et al. [18] showed the topological information, extracted from the relationships between sentences can be used in inference, can be applied to predict the very difficult legal entailment. Gholizadeh et al. [10] applied a different method for computing homological persistence to the task of authorship attribution, which is also a classification task.

These studies have demonstrated the usefulness of TDA in text understanding. However, none of them was for the task of contradiction detection, a fundamental but under-studied problem. Also importantly, none of them used TDA as an enhanced representation approach on top of deep learning representations.

3 METHOD

We adopted three widely used deep learning models as the base models: the continuous bag of words (CBOW), the Enhanced Sequential Inference Model (ESIM) from Chen et al.’s [3], which has the best documented accuracy in predicting contradicting pairs of sentences so far, and BERT (Bidirectional Encoder Representations from Transformers) [5], which has been proved as having the best performance in many text understanding and information retrieval tasks in recent three years. As for the implementation of TDA, we utilized Ripser.py [22], a convenient PH package for Python.

The topological feature we are concerned is a series of β_i s, each representing the number of i -th dimensional “holes”. Each hole at any dimension has the characteristics of birth time, death time, and persistence duration. Different sentence pairs may have different number of “holes” at each dimension, and therefore different series of β_i values. However, we need to specify a fixed dimension size for the topological feature representation because the size of input for a deep learning model is fixed. According to Pereira et al. [17], topological features with longer persistence duration are more important reflections on the characteristics of the data. Therefore

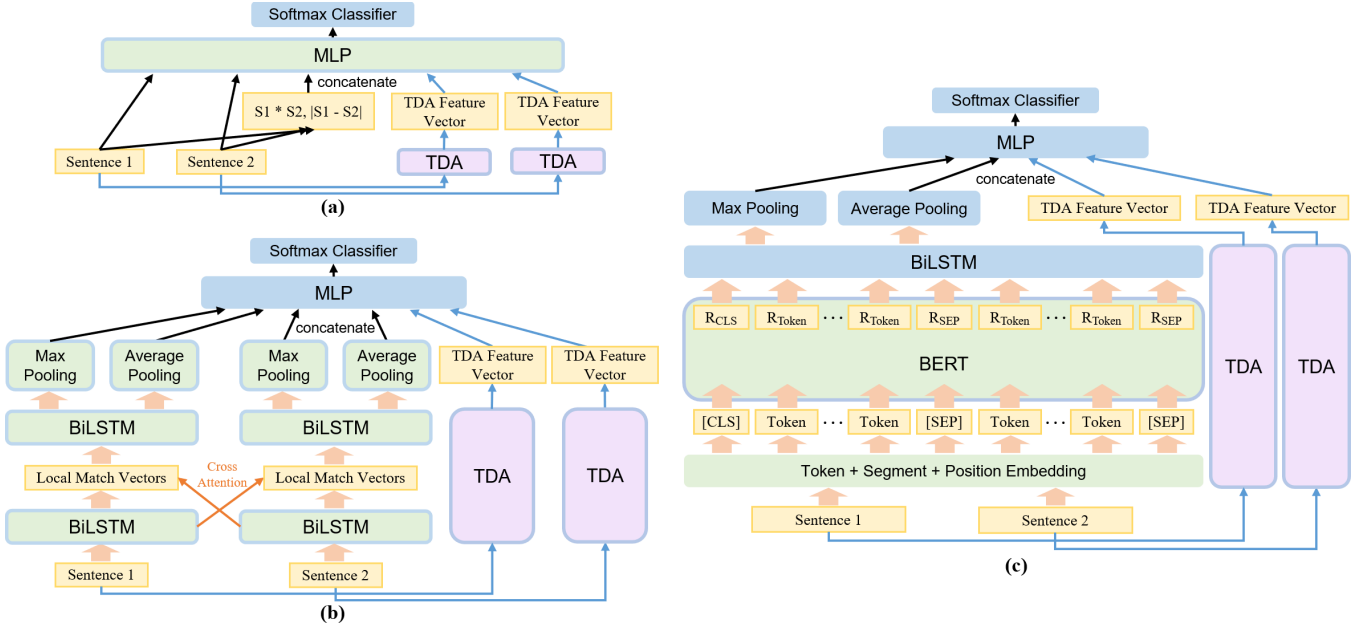


Figure 2: The Framework of Incorporating the TDA Representations into (a) CBOW, (b) ESIM, and (c) BERT

we selected the top k_i "holes" with the longest persistence duration for the i -th dimension as our feature input.

In addition, the number of "holes" decays quickly as the dimension increases. We only use the 0-dimensional and 1-dimensional topological features in our study, because according to our experiments, higher dimensional topological features do not exist in most sentence pairs. Therefore, we selected the top k_0 0-dimensional and the top k_1 1-dimensional topological features with longest persistence duration as the input of the model.

As for the incorporation of the TDA representation to deep learning models, we concatenated the TDA representation with the output of the final representation layer of each of the three base models, as shown in Figure 2. The concatenated representation was then the input for the classification layer to get the prediction. For the CBOW base model, such concatenation was with the pre-trained sentence embedding since the CBOW model does not have a representation layer. For ESIM, the concatenation was with the pooling representation of the output of the second BiLSTM, the last representation layer in the ESIM structure. For the BERT base model, we added a BiLSTM representation layer before the pooling of the word representations, which has demonstrated in our experiments to have a better model prediction performance than the vanilla BERT. Therefore, the concatenation was with the pooling representation of the output of this BiLSTM layer.

4 PROPOSED TYPES OF CONTRADICTIONS AND GROUND TRUTH DATA COLLECTION

We re-purposed the MultiNLI dataset [25], which contains approximately 433,000 pairs of sentences in 10 genres, such as government reports, fictions, telephone transcriptions, etc. The original

MultiNLI dataset has three possible labels for each record: *entailment*, *neutral*, and *contradiction*. We kept the original *contradiction* label and grouped the *entailment* and *neutral* labels into a new label *non-contradiction*. There are two labels in our new dataset: *contradiction* and *non-contradiction*.

After scanning a large amount of contradiction sentence pairs from the dataset, we observed that some contradictions were expressed in an explicit way using obvious indicator words, such as antonyms, negation, or numeric inconsistency, while others were more implicit. We hold meetings to examine the patterns and regularities of contradictions in sentence pairs, and independently suggested contradiction types. The process was iterative, reinforcing our existing understanding of contradiction types, while at the same time having maximum flexibility to be open to new types emerging from the data. As the result six types were proposed: **Negation**, using negative words (no, not, never, nobody, nothing, etc.) or phrases (neither...nor..., etc.) in a sentence to directly refute the meaning of the other; **Antonym**, using antonyms of the words or phrases in a sentence to refute the meaning of the other; **Replacement**, replacing a concept or an entity in a sentence to make it inconsistent with the other; **Switch**, changing or switching the place of some elements in a sentence; **Scope**, narrowing down or broadening up the scope expressed in a sentence; and **Latent**, refuting one sentence meaning without obvious indicator words, phrases, or structures. Table 1 shows the sentence pair examples for each of the types.

We invited three people fluent in English to label a randomly selected 1,000 contradictory sentence pairs from the 10 genres of MultiNLI according to our classification system. First, two people provided their initial labels independently. If there was agreement, the contradiction type label was assigned to that sentence pair. Otherwise, the third person was involved to break the tie. Each of

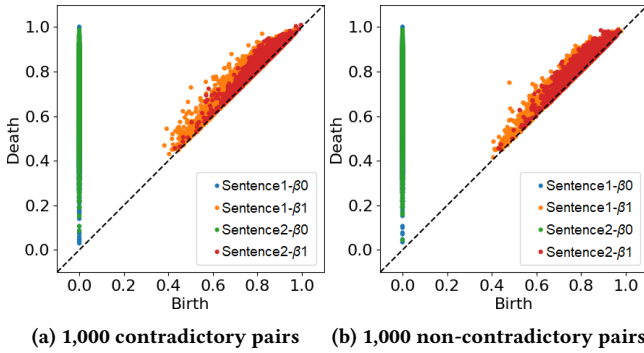
Table 1: Examples for Six Contradiction Types and Their Number of Cases in the 1,000 Contradiction Samples

Contradiction Type	Sentence 1	Sentence 2	Number of Cases
Negation	Amy and I both fought him and he nearly took us.	Neither Amy nor myself have ever fought him.	448
Antonym	Hills and mountains are sanctified in the cult of this religion.	The cult of that religion hates nature.	242
Replacement	The union has about 4,000 members in Canada.	There are 100 members in the union that live in Canada.	174
Switch	In the late 1990s, these extremist groups suffered major defeats by Kurdish forces.	Kurdish forces suffered devastating losses to these extremist groups.	10
Scope	Fun for adults and children.	Fun for only children.	12
Latent	Oh, Czarek, but will happen to us?	Czarek was not asked any questions on that day.	132

the three people was open-mined to new types if needed. As the result, 18 of the 1,000 sentence pairs have more than one type labels and the rest have one each. The type label distribution over the 1,000 sentence pairs is also presented in Table 1.

5 EXPERIMENTS

We have conducted experiments to evaluate the helpfulness of TDA representations in identifying contradictions.

**Figure 3: Visualizaiton of Birth and Death Time for "Holes"**

5.1 Experimental Setup

Table 2: Performance Comparisons of CBOW with Different Combinations of TDA Vectors

	Matched Test Set			Mismatched Test Set		
	Accuracy	Precision	Recall	Accuracy	Precision	Recall
$k_0=260, k_1=30$	0.753	0.723	0.399	0.761	0.756	0.409
$k_0=260, k_1=40$	0.753	0.723	0.404	0.761	0.753	0.395
$k_0=260, k_1=50$	0.753	0.733	0.389	0.762	0.764	0.403
$k_0=260, k_1=60$	0.753	0.722	0.405	0.761	0.758	0.412
$k_0=200, k_1=50$	0.753	0.724	0.401	0.761	0.759	0.410
$k_0=300, k_1=50$	0.753	0.730	0.382	0.761	0.764	0.397

Notes: All numbers are the average of five-fold cross validation results.

For the CBOW model, we followed the the structure in [25]. The concatenated representation was passed to a multilayer perceptron (MLP) with 3 layers, 300 nodes in each layer, and a 0.1 dropout rate in the last layer. The pre-trained embedding vectors were initialized with 300 dimensions. The out-of-vocabulary (OOV) words were

Table 3: Performance Comparisons of Different Models

Model	Matched Test Set			Mismatched Test Set		
	Accuracy	Precision	Recall	Accuracy	Precision	Recall
CBOW	0.753	0.718	0.404	0.761	0.752	0.414
CBOW+TDA	0.753	0.733	0.389	0.762	0.764	0.403
ESIM	0.822	0.777	0.641	0.828	0.806	0.630
ESIM+TDA	0.826	0.791	0.637	0.833	0.817	0.637
BERT	0.851	0.803	0.725	0.859	0.820	0.733
BERT+TDA	0.854	0.821	0.710	0.859	0.838	0.711

Notes: All numbers are the average of five-fold cross validation results.

Table 4: Model Accuracy on Different Contradiction Types

Model	Contradiction Type					
	Negation	Antonym	Replacement	Switch	Scope	Latent
CBOW	0.622	0.297	0.086	0.100	0.333	0.333
CBOW+TDA	0.612	0.285	0.086	0.100	0.167	0.364
ESIM	0.850	0.632	0.218	0.200	0.583	0.598
ESIM+TDA	0.877	0.628	0.253	0.200	0.583	0.621
BERT	0.897	0.686	0.448	0.100	0.667	0.644
BERT+TDA	0.891	0.719	0.460	0.200	0.750	0.682

Notes: All numbers are the average of five-fold cross validation results.

all set to 0. We truncated both sentences at the length of 25. For ESIM, each BiLSTM layer contained 300 LSTM blocks. The sentence input was the pre-trained embedding representation with the same parameters as the input of the CBOW base model. The setup of the MLP layer was also the same with the CBOW base model. For BERT, we used the BERT-base model. The added BiLSTM layer contained 64 LSTM blocks. Sentence length was set to 128.

The new dataset was separate into the training and test datasets in the same way with the original MultiNLI dataset. The data in the training set belonged to the first five out of the ten genres. The test set is further divided into two subsets: one is called the Matched Test Set whose genres are the same as the training set, and the other is called Mismatched Test Set whose genres are the last five genres outside the training data genres.

5.2 Model Results on Contradiction Detection

For the 1,000 selected contradictory sentence pairs and another 1,000 randomly selected non-contradictory sentence pairs, we have visualized their PH - birth time and death time of the β_0 and β_1 "holes" in a two-dimensional coordinate. As in Figure 3, different

Table 5: Performance Comparisons of Different Models on Different Genres

Model	Genre														
	slate			government			telephone			travel			fiction		
	Accuracy	Precision	Recall	Accuracy	Precision	Recall	Accuracy	Precision	Recall	Accuracy	Precision	Recall	Accuracy	Precision	Recall
CBOW	0.720	0.623	0.348	0.767	0.738	0.433	0.761	0.754	0.416	0.740	0.701	0.417	0.763	0.778	0.398
CBOW+TDA	0.728	0.640	0.336	0.767	0.745	0.420	0.761	0.770	0.398	0.750	0.728	0.397	0.761	0.778	0.388
ESIM	0.794	0.713	0.590	0.843	0.825	0.653	0.835	0.807	0.662	0.809	0.756	0.630	0.829	0.785	0.667
ESIM+TDA	0.799	0.730	0.588	0.850	0.837	0.665	0.835	0.805	0.663	0.813	0.781	0.610	0.832	0.799	0.658
BERT	0.827	0.750	0.686	0.873	0.846	0.744	0.856	0.815	0.732	0.846	0.797	0.722	0.856	0.809	0.739
BERT+TDA	0.831	0.774	0.662	0.877	0.860	0.742	0.859	0.830	0.721	0.850	0.822	0.700	0.854	0.816	0.722

Model	letters			verbatim			facetoface			oup			nineeleven		
	Accuracy	Precision	Recall	Accuracy	Precision	Recall	Accuracy	Precision	Recall	Accuracy	Precision	Recall	Accuracy	Precision	Recall
	Accuracy	Precision	Recall	Accuracy	Precision	Recall	Accuracy	Precision	Recall	Accuracy	Precision	Recall	Accuracy	Precision	Recall
CBOW	0.793	0.797	0.501	0.741	0.701	0.384	0.762	0.777	0.400	0.763	0.784	0.37	0.749	0.705	0.416
CBOW+TDA	0.780	0.774	0.456	0.748	0.720	0.386	0.762	0.786	0.388	0.759	0.778	0.363	0.755	0.733	0.409
ESIM	0.855	0.836	0.694	0.807	0.770	0.593	0.828	0.805	0.639	0.820	0.810	0.59	0.818	0.781	0.627
ESIM+TDA	0.863	0.853	0.701	0.814	0.789	0.598	0.832	0.815	0.642	0.828	0.819	0.609	0.821	0.786	0.630
BERT	0.892	0.868	0.793	0.837	0.788	0.696	0.855	0.807	0.744	0.858	0.830	0.705	0.850	0.805	0.723
BERT+TDA	0.889	0.876	0.772	0.838	0.817	0.658	0.856	0.821	0.724	0.860	0.846	0.695	0.854	0.827	0.706

Notes: The genre descriptions are available in [25]. All numbers are the average of five-fold cross validation results.

topological patterns could be observed for the 1,000 contradictory pairs and the 1,000 non-contradictory pairs. For example, the distribution difference of the β_1 points for Sentence1 and Sentence2 looks larger and more disperse for Figure 3(a) (the contradictory pairs) than Figure 3(b) (the non-contradictory pairs).

Since the embedding vector of each sentence has the dimensionality of 300, we set the size of TDA vectors ($k_0 + k_1$) to around 300 to prevent TDA from having too much or too little effect on the model. We started with $k_0 = 260$ and $k_1 = 40$ as the number of β_0 is usually an order of magnitude larger than that of β_1 . We experimented with different combinations of values of k_0 and k_1 . The model performance is shown in Table 2. We found that $k_0 = 260$ and $k_1 = 50$ were the best.

The performance of the three models on the Matched and Mismatched test sets is summarized in Table 3. Overall speaking, the help of incorporating the TDA representations for ESIM and BERT is larger than that for CBOW according to the pairwise comparison of the same model with and without TDA. This suggests that TDA representations are more useful when the base model has a representation layer. The concatenated representation is able to improve the performance of contradiction detection. We do not observe any clear model performance difference between matched and mismatched testing datasets, suggesting that all the trained models are able to make predictions beyond the text genres seen in the training datasets. From the model perspective, the BERT model, especially with the TDA representations, has the best performance in terms of accuracy and recall at a little cost of recall, meaning it is more precise but a little less sensitive. According to De Marnette et al. [4], since contradiction is relatively rare compared to non-contradiction in real-world text, precision is more valuable than recall in the task of contradiction detection. We have seen the improvement of precision with the help of TDA for all of the three models.

Table 4 shows the model accuracy for each of the six contradiction types. Overall, the BERT model with the TDA representations achieved the highest accuracy for all the contradiction types except **Negation**, demonstrating the helpfulness of topological features in finding contradictions, especially the **Latent** contradictions, and

those complicated types, such as **Replacement** and **Switch**, which do not rely on changing words but a sentence structure to express contradictions. Generally speaking, **Negation** and **Antonym** are the least challenging contradictions to detect whereas the other four types, **Replacement**, **Switch**, **Scope**, and **Latent** are the challenging ones with relatively lower accuracy for all the models.

We are also interested in seeing if the performance of the three models (CBOW, ESIM, and BERT), with or without the TDA representations, varies on different text genres. Table 5 shows the results. We have seen the improvement for each model with the TDA representation in almost all the genres in terms of accuracy and precision, with a little sacrifice of recall, a similar trend with the results in Table 3. The performance among different genres do not have differences, meaning the models were appropriately trained to have a good understanding of the general language regardless of written or oral English, fiction or non fiction, etc.

6 CONCLUSION

In this paper, we proposed a framework of incorporating TDA representations into three widely used deep learning models to identify contradictions in sentence pairs. In addition, six types of contradictions were proposed and a testing dataset with human labeled contradiction types was collected. We evaluated the performance of the deep learning models with the TDA representations, as well as their variations on different contradiction types and text genres. Positive and promising results have been obtained for the usefulness of the TDA representations on top of deep learning representations.

ACKNOWLEDGMENTS

This research is supported by National Science Foundation (NSF) (Award #1910696). We would like to thank NSF to make this research possible.

REFERENCES

- [1] Fakhri Abbas and Xi Niu. 2019. One size does not fit all: Modeling users' personal curiosity in recommender systems. *ArXiv.org* (2019).

- [2] Gunnar Carlsson. 2009. Topology and data. *Bull. Amer. Math. Soc.* 46, 2 (2009), 255–308.
- [3] Qian Chen, Xiaodan Zhu, Zhenhua Ling, Si Wei, Hui Jiang, and Diana Inkpen. 2016. Enhanced LSTM for natural language inference. *arXiv preprint arXiv:1609.06038* (2016).
- [4] Marie-Catherine De Marneffe, Anna N Rafferty, and Christopher D Manning. 2008. Finding contradictions in text. In *Proceedings of ACL-08: HLT*. 1039–1047.
- [5] Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2018. Bert: Pre-training of deep bidirectional transformers for language understanding. *arXiv preprint arXiv:1810.04805* (2018).
- [6] Pratik Doshi and Wlodek Zadrozny. 2018. Movie genre detection using topological data analysis. In *International Conference on Statistical Language and Speech Processing*. Springer, 117–128.
- [7] Lorik Dumani, Patrick J Neumann, and Ralf Schenkel. 2020. A framework for argument retrieval. In *European Conference on Information Retrieval*. Springer, 431–445.
- [8] Herbert Edelsbrunner and John Harer. 2010. *Computational topology: an introduction*. American Mathematical Soc.
- [9] Herbert Edelsbrunner, David Letscher, and Afra Zomorodian. 2000. Topological persistence and simplification. In *Proceedings 41st annual symposium on foundations of computer science*. IEEE, 454–463.
- [10] Shafie Gholizadeh, Armin Seyeditabari, and Wlodek Zadrozny. 2018. Topological signature of 19th century novelists: Persistent homology in text mining. *Big Data and Cognitive Computing* 2, 4 (2018), 33.
- [11] He-Liang Huang, Xi-Lin Wang, Peter P Rohde, Yi-Han Luo, You-Wei Zhao, Chang Liu, Li Li, Nai-Le Liu, Chao-Yang Lu, and Jian-Wei Pan. 2018. Demonstration of topological data analysis on a quantum processor. *Optica* 5, 2 (2018), 193–198.
- [12] Chuqin Li, Xi Niu, Ahmad Al-Doulat, and Noseong Park. 2018. A computational approach to finding contradictions in user opinionated text. In *2018 IEEE/ACM International Conference on Advances in Social Networks Analysis and Mining (ASONAM)*. IEEE, 351–356.
- [13] Pek Y Lum, Gurjeet Singh, Alan Lehman, Tigran Ishkanov, Mikael Vejdemo-Johansson, Muthu Alagappan, John Carlsson, and Gunnar Carlsson. 2013. Extracting insights from the shape of complex data using topology. *Scientific reports* 3, 1 (2013), 1–8.
- [14] Xi Niu and Fakhri Abbas. 2019. Computational surprise, perceptual surprise, and personal background in text understanding. In *Proceedings of the 2019 Conference on Human Information Interaction and Retrieval*. 343–347.
- [15] Xi Niu, Fakhri Abbas, Mary Lou Maher, and Kazjon Grace. 2018. Surprise me if you can: Serendipity in health information. In *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems*. 1–12.
- [16] Xi Niu, Wlodek Zadrozny, Kazjon Grace, and Weimao Ke. 2018. Computational surprise in information retrieval. In *The 41st International ACM SIGIR Conference on Research & Development in Information Retrieval*. 1427–1429.
- [17] Cássio MM Pereira and Rodrigo F de Mello. 2015. Persistent homology for time series and spatial data clustering. *Expert Systems with Applications* 42, 15-16 (2015), 6026–6038.
- [18] Ketki Savle, Wlodek Zadrozny, and Minwoo Lee. 2019. Topological data analysis for discourse semantics?. In *Proceedings of the 13th International Conference on Computational Semantics-Student Papers*. 34–43.
- [19] Donald R Sheehy. 2013. Linear-size approximations to the Vietoris–Rips filtration. *Discrete & Computational Geometry* 49, 4 (2013), 778–796.
- [20] Christian Stab, Johannes Daxenberger, Chris Stahlhut, Tristan Miller, Benjamin Schiller, Christopher Tauchmann, Steffen Eger, and Iryna Gurevych. 2018. Argumenttext: Searching for arguments in heterogeneous sources. In *Proceedings of the 2018 conference of the North American chapter of the association for computational linguistics: demonstrations*. 21–25.
- [21] Jianlin Su, Jiarun Cao, Weijie Liu, and Yangyiwen Ou. 2021. Whitening sentence representations for better semantics and faster retrieval. *arXiv preprint arXiv:2103.15316* (2021).
- [22] Christopher Tralie, Nathaniel Saul, and Rann Bar-On. 2018. Ripser. py: A lean persistent homology library for python. *Journal of Open Source Software* 3, 29 (2018), 925.
- [23] Henning Wachsmuth, Martin Potthast, Khalid Al Khatib, Yamen Ajjour, Jana Puschmann, Jiani Qu, Jonas Dorsch, Viorel Morari, Janek Bevendorff, and Benno Stein. 2017. Building an argument search engine for the web. In *Proceedings of the 4th Workshop on Argument Mining*. 49–59.
- [24] Henning Wachsmuth, Shahbaz Syed, and Benno Stein. 2018. Retrieval of the best counterargument without prior topic knowledge. In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*. 241–251.
- [25] Adina Williams, Nikita Nangia, and Samuel R Bowman. 2017. A broad-coverage challenge corpus for sentence understanding through inference. *arXiv preprint arXiv:1704.05426* (2017).
- [26] Xiaojin Zhu. 2013. Persistent Homology: An Introduction and a New Text Representation for Natural Language Processing.. In *IJCAI*. 1953–1959.