

REFeree REPORT FOR
“TORSION IN PERSISTENT HOMOLOGY AND NEURAL NETWORKS”
BY MARIA WALCH

1. HIGH-LEVEL SUMMARY

If I had to isolate one big idea from this paper, it is approximately this:

We can view an encoder as a *continuous*¹ map $f_{enc} : \mathbb{R}^d \rightarrow \mathbb{R}^\lambda$ where $\lambda \ll d$, regardless of how it arises as layers in a neural network. If $X \subset \mathbb{R}^d$ is a point cloud, viewed as an input data set to the encoder, then if $L := f_{enc}(X) \subseteq \mathbb{R}^\lambda$ is the latent representation of the point cloud, then Corollary 3.4.6 of [1] implies three things:

- (1) $H_i(L; \mathbb{Z}) = 0$ for $i \geq \lambda$,
- (2) $H_i(L; \mathbb{Z})$ is torsion free for $i = \lambda - 1$, and finally
- (3) $H_i(L; \mathbb{Z})$ is torsion free for $i = \lambda - 2$.

Since an autoencoder is a composition of an encoder followed by a decoder

$$f_{auto} = f_{dec} \circ f_{enc}$$

and since homology is functorial, then torsion in X in homological degrees $i \geq \lambda - 2$ will automatically map to zero through the autoencoder. In other words, **a topological autoencoder cannot represent torsion in degrees $i \geq \lambda - 2$.**

Of course, one critical omission in the above summary is that we need to talk about *persistent* homology, as $L := f(X)$ is currently the image of a finite point cloud, which has no interesting homology on its own. **One critical mistake the author makes is that they don’t phrase their results in terms of the alpha filtration of the point cloud X , which is necessary for viewing persistence as a summary of some *embedded* object.** It is easy to miss the alpha filtration assumption in [2], but it is there. The author currently states things in terms of the Rips filtration, which is wrong. As well documented by Henry Adams and others, the Rips filtration can have homology in degrees higher than the embedding space, and Alexander duality (the workhorse of Corollary 3.4.6 of [1]) does not apply to the Rips filtration.

However, even if the author were to correct this—easy to fix!—mistake, there is still a broader question of whether it is worth publishing such a result in *JACT*. For the record, I have seen less sophisticated topological results [3], when presented well and connected to explicit neural network algorithms such as message passing, get published in prestigious conference workshops. My feeling is that if the author wrote a clear paper, stating the “big idea” above correctly, and performing some numerical experiments, it would be good conference publication. However, the paper under review substitutes clear and easy to state mathematical results with needlessly confusing wrong statements, e.g., Propositions 3, 4, 6 and 7, as well as Corollaries 3, 4, and 5. This is the bulk of the theoretical work and it cannot be salvaged with easy fixes. **As such I recommend rejection.**

Overall, the author makes several statements where the hypotheses contradict themselves and the proofs do not make sense. Terms like “preserves torsional structure” and “by the representability

¹This should be emphasized as homology is only functorial with respect to continuous maps.

of torsion in the latent space” are not precisely defined anywhere. Moreover, concepts from [4] are recalled incorrectly and appear to be incorrectly applied. More detailed comments follow.

2. DETAILED COMMENTS

- (1) Section 2.2. First, I would restate Theorem 1 with “for all” logical quantifiers rather than “for any”. I realize this is how it appears in [2], but I think for native English speakers it makes it more explicit that we need the condition for all pairs of indices $i < j$, for example.
- (2) Equation (4), the subscript nn in $x_{i,nn}$ should be explained.
- (3) Equation (5) makes it seem like W_{jk} is a d -vector, but at layer j , the dimension of the image of x_i through the previous $j - 1$ layers certainly has changed its dimension. It should be made clear that the dimension of the input to layer j may not be the same dimension as the input $\mathcal{D}$ to the whole network.
- (4) Assumption 1: This is not an assumption. This is just a statement that the output of the j^{th} layer is a vector. Calling this “Euclidicity” and identifying this as an assumption doesn’t make sense to me.
- (5) In Equation (6), the subscript l in $x_{i,l}$ should be explained.
- (6) First paragraph after Equation (9) includes the statement “If the encoder does not exhibit stochasticity, the encoder is surjective.” This doesn’t make any sense to me. I think you might mean “Generic encoders are surjective.”? If so, a proof or reference should be given. I can see that *linear* maps $\mathbb{R}^d$ to $\mathbb{R}^\lambda$ with $\lambda \ll d$ are generically surjective (= full rank), but I suspect certain choices of activation functions would remove surjectivity (like for sigmoid activation functions).
- (7) Example 1: “relataive” is a misspelling of “relative.”
- (8) In Section 5.1, when recalling the work of [2], make sure to remind the reader that they are looking at the alpha filtration.
- (9) Top of page 10: It would be helpful to briefly summarize the results from Appendix A—ReLU and Leaky ReLU lose torsion, but the other activation functions do not.
- (10) Start of Section 6: I think the basic idea here is wrong. You assert that if there is torsion in a high-dimensional homology group of X (or some simplicial complex associated to X), then an encoder must be “highly representative”—you mean faithfully represented, i.e., the same homology—in lower dimensions. This is simply impossible, for obvious reasons.
- (11) Proposition 3: Define “preserves torsional structure”. If what you mean is that f maps torsion elements to torsion elements, then this is always true. To wit, if $f : X \rightarrow L$ is a continuous map, then $f_* : H_*(X; \mathbb{Z}) \rightarrow H_*(L; \mathbb{Z})$ is a group homomorphism. In particular if $\alpha \in H_*(X; \mathbb{Z})$ is torsion, i.e., there exists $q^n \alpha = 0$, then we know that $f_*(q^n \alpha) = q^n f_*(\alpha) = 0$, i.e., $f_*(\alpha)$ is torsion. However, as you already observed $L \subseteq \mathbb{R}^\lambda$, so when $p \geq \lambda - 2$ we know by Corollary 3.4.6 of [1] that $H_p(L; \mathbb{Z})$ is torsion free, hence $f_*(\alpha) = 0$. I really don’t know why you’re introducing another prime q' into the discussion. If you’re trying to say that f doesn’t induce an isomorphism on homology, then I agree, and it’s true with $q' = q$, but this just reduces to Proposition 2 and this new proposition offers nothing new.
- (12) Proof of Proposition 3: I’m confused from the first sentence already. “The appearance of torsional structures in lower-dimensional homology implies lower-dimensional (relative) homology groups ($p < \lambda - 1$), which lacked torsion in the input data, now exhibit torsion in the latent space.” First of all, the case $p < \lambda - 1$ is not in the scope of discussion. Prop 3 is talking about $p \geq \lambda - 1$. It appears you’re adding a further assumption about X , namely

that it has no torsion in homological degrees $p \leq \lambda - 1$. OK. Now how did torsion appear in the latent space? It's possible that $H_p(X; \mathbb{Z})$ is torsion free, but $H_p(L; \mathbb{Z})$ may or may not have torsion. What hypotheses cause L to “exhibit torsion”? Again if all you're saying is that an encoder might not preserve integral homology, I'm convinced, but it has nothing to do with inducing q' torsion somewhere else. Looking at the statement of the proposition, where $p \geq \lambda - 1$, the only interesting case is $p = \lambda - 1$. We know that $H_{p=\lambda-1}(L; \mathbb{Z})$ is torsion free and I've proved above that $f_*(\alpha) = 0$, necessarily, so “preserving torsional structures” is impossible in this degree. Consequently, we cannot induce isomorphisms in integral homology, which is just Prop 2. Are you just trying to say that, because of Cor 3A.7 of [1], we can also detect this lack of isomorphism by considering field coefficients over $\mathbb{F}_{q'}$ for some different prime q' ? If so, just say that!

- (13) Prop 4: Again “The representability of torsion in the latent space” is not a theorem, or an axiom, from what I can tell, so how can you derive a statement from this. I think this statement that torsion in homological degree $p \geq \lambda - 1$ for X automatically implies non-trivial torsion in two homological degrees in L smaller than $\lambda - 1$ is just false. If not, please give at least one example where this phenomenon occurs.
- (14) Proof of Prop 4: Your proof oscillates from saying obvious general things—“elements of the torsion subgroup lie in the kernel of boundary operator” (all homology classes are by definition in the kernel of the boundary operator)—to oddly specific wrong things—“cycles whose torsion coefficient times their representative loop bounds a simplex...are incorrectly identified as non-trivial homology classes over $\mathbb{Z}/q\mathbb{Z}$ ”. First of all to say α is the representative of a torsion class $[\alpha] \in H_p$ is to exactly say that $q^n \alpha$ is a boundary of some other chain, not necessarily a simplex. Assuming $[\alpha] \neq 0$ over $\mathbb{Z}$ coefficients means its rightfully non-trivial over $\mathbb{Z}/q\mathbb{Z}$. It's not incorrectly identified.
- (15) Proof of Prop 4, Regarding Euler Characteristic: I think I see what you're doing here, but you're just trying to recast a classical observation into something about neural nets, which I don't think can be true. Again, if X has torsion in degree $p \geq \lambda - 1$, then there is no way that an encoder can preserve this information. In particular, there is no reason to think $\chi(X) = \chi(L)$. That's the starting point and is a missing assumption for what follows. The invariance of Euler characteristic, i.e., that for a cell complex X , decomposed as cells $\{\sigma_i^{n_i}\}$, we have the equalities

$$\chi(X) = \sum_i (-1)^{n_i} = \sum_p \text{rank} H_p(X; \mathbb{Z}) = \sum_p \dim H_p(X; \mathbb{Q}) = \sum_p \dim H_p(X; \mathbb{F}_q).$$

This last equality is a bit more tricky, but [5] proves it. Anyhow, if X has q -torsion in homological degree p , this is going to cause the dimension of $H_p(X; \mathbb{F}_q)$ to be higher than the rank of $H_p(X; \mathbb{Z})$. This seems to create a problem, but actually Corollary 3A.6(b) of [1] tells you exactly what happens: it will be compensated by an extra copy of $\mathbb{F}_q$ in $H_{p+1}(X; \mathbb{F}_q)$. This is what causes this “torsional feature” to cancel in pairs in the Euler characteristic calculation over the field $\mathbb{F}_q$. Notice the “extra feature” actually appears in a *higher* homological degree, rather than a lower.

- (16) The persistence entropy discussion: At this point there were too many red flags for me to continue reviewing the manuscript in detail. Lemma 1 incorrectly recalls Theorem 13 of [4]. The radius should be T and the diameter is $2T$.
- (17) Corollary 3 doesn't make sense to me. How are we supposed to know “that this bar represents a topological feature”? Don't all bars represent topological features? Isn't it a question of significance? I don't think a single detector of significance exists or can exist.

REFERENCES

- [1] “*Algebraic Topology*” by Allen Hatcher. Cambridge University Press, 2002. <https://pi.math.cornell.edu/~hatcher/AT/AT.pdf>
- [2] “*Field Choice Problem in Persistent Homology*” by Ippei Obayashi and Michio Yoshiwaki. Published in Discrete and Computational Geometry, August 2023. <https://link.springer.com/article/10.1007/s00454-023-00544-7>
- [3] “*Topological Blindspots: Understanding and Extending Topological Deep Learning Through the Lens of Expressivity*” by Yam Eitan, Yoav Gelberg, Guy Bar-Shalom, Fabrizio Frasca, Michael M. Bronstein, Haggai Maron. Published in the NeurIPS 2024 Workshop on Symmetry and Geometry in Neural Representations. <https://openreview.net/forum?id=ZtI66QxTYn>
- [4] “*Persistent entropy for separating topological features from noise in vietoris-rips complexes*” by Nieves Atienza, Rocio Gonzalez-Diaz, and Matteo Rucco. Published in the Journal of Intelligent Information Systems, 2017. <https://link.springer.com/article/10.1007/s10844-017-0473-4>
- [5] “*Euler characteristic: dependence on coefficients*” Math StackExchange answer by Najib Idrissi <https://math.stackexchange.com/users/10014/najib-idrissi> on 2015-07-11 in response to <https://math.stackexchange.com/q/1357181>.