

We need to talk about random seeds

Anonymous ACL submission

Abstract

Modern neural network libraries all take as a hyperparameter a random seed, typically used to determine the initial state of the model parameters. In this position piece, I argue that there are some appropriate uses for random seeds: as part of the hyperparameter search to select a good model, creating an ensemble of several variants of a model, or measuring the sensitivity of the training algorithm to the random seed hyperparameter. I argue against some inappropriate uses for random seeds: using a fixed random seed for “replicability” and creating score distributions for performance comparison. I review 85 recent publications from the ACL Anthology and find that more than 50% are using random seeds inappropriately.

1 Introduction

Modern neural network libraries all take as a hyperparameter a *random seed*, a number that is used to initialize a pseudorandom number generator. That generator is typically used to determine the initial state of neural network parameters, but may also be used for other purposes, such as shuffling the training data. Like any hyperparameter, neural network random seeds can have a large or small impact on model performance depending on the specifics of the model and the data.

The pre-trained transformer-based models that are currently popular in NLP (BERT, [Devlin et al., 2019](#); RoBERTa [Liu et al., 2019](#); etc.) have been observed to be quite sensitive to their random seeds ([Risch and Krestel, 2020](#); [Dodge et al., 2020](#); [Mosbach et al., 2021](#)). Several solutions to this problem have been proposed, including specific optimizer setups ([Mosbach et al., 2021](#)), ensemble methods ([Risch and Krestel, 2020](#)), and explicitly tuning the random seed like other hyperparameters ([Dodge et al., 2020](#)).

However, while it seems that the NLP community is reasonably conscious of the problems that

random seeds present, it is inconsistent in its approaches to solving those problems. In the remainder of this position piece, I first present a taxonomy of different ways that neural network random seeds are used in the NLP community, explaining which uses are appropriate and which are inappropriate. Then I review 85 articles recently published in the ACL Anthology, categorizing their random seed uses based on the taxonomy. I find that more than 50% of the articles use random seeds inappropriately, suggesting that the NLP community still needs a broader discussion about how we approach random seeds.

2 A taxonomy of random seed uses

In this section I highlight five common uses of neural network random seeds in the NLP community, and categorize them as either appropriate or inappropriate.

2.1 Appropriate uses

Hyperparameter selection The random seed is a hyperparameter of the neural network that determines where in the model’s parameter space optimization should begin. As the random seed is a hyperparameter, it can and should be tuned just as other hyperparameters are. Unlike some other hyperparameters, there is no intuitive explanation of why one random seed would be better or worse than another, so the typical strategy is to just try a number of randomly selected seeds. For example:

Instead, we compensate for the inherent randomness of the network by training multiple models with randomized initializations and use as the final model the one which achieved the best performance on the validation set... ([Björne and Salakoski, 2018](#))

The test results are derived from the 1-best random seed on the validation set.

080	(Kuncoro et al., 2020)	For consistency, we used the same set	124
		of hyper-parameters and a fixed random	125
081	Ensemble creation Ensemble methods are an	seed across all experiments. (Lin et al.,	126
082	effective way of combining multiple machine-	2020)	127
083	learning models to make better predictions		
084	(Rokach, 2010). A common approach to creating	Why is this inappropriate? First, fixing the random	128
085	neural network ensembles is to allow multiple train-	seed does not guarantee replicability. For example,	129
086	ing runs of the same architecture, each starting with	the tensorflow library has a history of producing	130
087	a different random seed, to vote in the ensemble	different results even given the same random seeds,	131
088	(Perrone and Cooper, 1995). For example:	especially on GPUs (Two Sigma, 2017; Kanwar	132
		et al., 2021). Second, not tuning the random seed	133
089	<i>In order to improve the stability of the</i>	hyperparameter has the same drawbacks as not tun-	134
090	<i>RNNs, we ensemble five distinct models,</i>	ing any other hyperparameter: performance will	135
091	<i>each initialized with a different random</i>	be an underestimate of the performance of a tuned	136
092	<i>seed. (Nicolai et al., 2017)</i>	model.	137
		Instead, the random seed should be tuned as	138
093	<i>Our model is composed of the ensem-</i>	any other hyperparameter. Dodge et al. (2020),	139
094	<i>ble of 8 single models. The hyper-</i>	for example, show that doing so leads to simpler	140
095	<i>parameters and the training procedure</i>	models exceeding the published results of more	141
096	<i>used in each single model are the same</i>	complex state-of-the-art models on multiple GLUE	142
097	<i>except the random seed. (Yang and</i>	tasks (Wang et al., 2018). For transparency, when	143
098	<i>Wang, 2019)</i>	the random seed is tuned in this way, both the space	144
		of random seeds searched and the final selected	145
099	Sensitivity analysis Sometimes it is useful to	random seed should be reported.	146
100	demonstrate how sensitive a neural network is to a		
101	particular hyperparameter. For example, Santurkar	Performance comparison It is a good idea to	147
102	et al. (2018) shows that batch normalization makes	compare not just point estimates of a model’s per-	148
103	neural networks less sensitive to the learning rate	formance, but distributions of model performance,	149
104	hyperparameter. Similarly, it may be useful to show	as comparing performance distributions may result	150
105	how sensitive neural networks are to their random	in more reliable conclusions. It has been suggested	151
106	seed hyperparameter. For example:	that such distributions can be obtained by running	152
		the same model with different random seeds. For	153
107	<i>We next (§3.3) examine the expected vari-</i>	example:	154
108	<i>ance in attention-produced weights by</i>		
109	<i>initializing multiple training sequences</i>	<i>Instead of publishing and reporting sin-</i>	155
110	<i>with different random seeds... (Wiegr-</i>	<i>gle performance scores, we propose</i>	156
111	<i>effe and Pinter, 2019)</i>	<i>to compare score distributions based</i>	157
		<i>on multiple executions. (Reimers and</i>	158
112	<i>Our model shows a lower standard de-</i>	<i>Gurevych, 2017)</i>	159
113	<i>viation on each task, which means our</i>		
114	<i>model is less sensitive to random seeds</i>	<i>Indeed, the best approach is to stop re-</i>	160
115	<i>than other models. (Hua et al., 2021)</i>	<i>porting single-value results, and instead</i>	161
		<i>report the distribution of results from a</i>	162
116	2.2 Inappropriate uses	<i>range of seeds. (Crane, 2018)</i>	163
117	Single fixed seed NLP articles sometimes pick	Why is this inappropriate? Running the same	164
118	a single fixed random seed, claiming that this is	model with multiple random seeds is a good way	165
119	done to improve consistency or replicability. For	to estimate the sensitivity of the model to the ran-	166
120	example:	dom seed hyperparameter. But that’s not what one	167
		is looking for when comparing models. The au-	168
121	<i>An arbitrary but fixed random seed was</i>	thors above would probably not suggest this for	169
122	<i>used for each run to ensure reproducibil-</i>	any other hyperparameter. Should we generate a	170
123	<i>ity... (Le and Fokkens, 2018)</i>	distribution over model performance by training	171

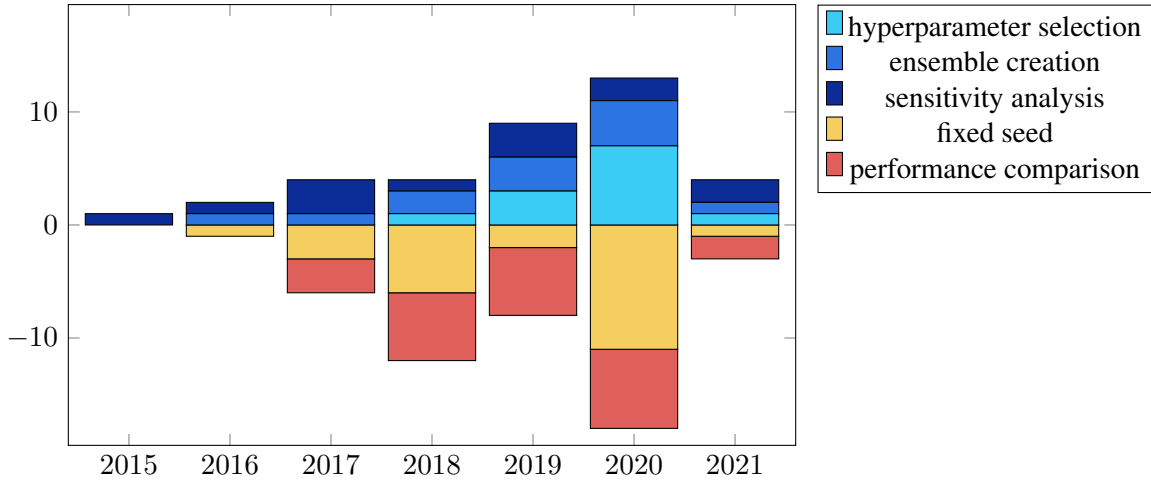


Figure 1: Uses of neural network random seeds by year for 85 ACL Anthology articles.

the same model but with different learning rates? No, that would be generating a bunch of suboptimal models within the distribution we’re trying to compare. For the same reason, we don’t want a bunch of models that are suboptimal in their choice of random seed.

Instead, to generate distributions over model performance, we should use standard statistical techniques for doing this. For example, bootstrap samples may be drawn from the test set, and evaluating a model on each of those samples will give a distribution over the model’s expected performance (Dror et al., 2018). As long as the models applied to these bootstrap samples have been tuned appropriately (including tuning the random seed), comparing two models based on such bootstrap distributions will allow a statistically sound assessment of the performance difference.

3 State of random seed uses in ACL

Having introduced both appropriate and inappropriate uses of neural network random seeds, I now turn to the current state of NLP with respect to such seeds.

On 29 Jun 2021, I searched the ACL Anthology for articles containing the phrases “random seed” and “neural network”¹. The ACL Anthology search interface returns a maximum of 10 pages of results, with 10 results per page, so I collected 100 search results. Non-articles (entire proceedings, author pages, supplementary material) were excluded, as were articles where the random seeds were not

¹<https://www.aclweb.org/anthology/search/?q=%22random+seed%22+%22neural+network%22>

Type	Purpose	Count
Appropriate	Hyperparameter selection	12
Appropriate	Ensemble creation	13
Appropriate	Sensitivity analysis	12
Appropriate sub-total		37
Inappropriate	Fixed seed	24
Inappropriate	Performance comparison	24
Inappropriate sub-total		48

Table 1: Uses of neural network random seeds for 85 ACL Anthology articles.

used to initialize a neural network (e.g., they were used only for dataset selection). The result was 85 articles, from publications between 2015 and 2021.

I read each of the articles and categorized its use of random seeds into one of the five purposes introduced in section 2. The supplementary material for this article includes a spreadsheet detailing each article reviewed, its purpose for using neural network random seeds, and a snippet of text from the article justifying my assignment of that purpose category.

Table 1 shows the distribution of articles across the different random seed purposes. More than half of the articles (48) include an inappropriate use of random seeds, with fixed seeds and performance comparisons being equally likely. This suggests that NLP researchers are frequently misusing neural network random seeds.

One might wonder if the NLP community is getting better over time, that is, if inappropriate uses are on the decline as NLP researchers become more familiar with neural networks research. Figure 1

shows that this is not the case: though the volume of articles that matched the query varies from year to year, for most years the number of inappropriate uses of random seeds is similar to the number of appropriate uses. This suggests that NLP researchers continue to have trouble distinguishing appropriate from inappropriate uses of neural network random seeds.

4 Discussion

We have seen that inappropriate uses of neural network random seeds – including using only a fixed seed or using random seeds to generate performance distributions for model comparisons – are still widespread within the NLP community. The analysis here is probably a conservative estimate of the problem. Articles only matched the query if they had the explicit phrases “neural network” and “random seed” both within the article. That means the search did not return articles on neural networks where no “random seed” was mentioned, yet in such cases it is likely that a single fixed seed was used. Therefore the proportion of fixed seed papers in our sample is likely an underestimate of the proportion in the true population.

How do we move the NLP community toward more appropriate uses of neural network random seeds? I hope that this article can help to start the necessary conversations, but clearly it is not an endpoint in and of itself. Part of the responsibility must fall on mentors in the NLP community, such as university faculty and industry research leads, to ensure that they are training their mentees about these topics. Part of the responsibility will fall on reviewers of NLP articles, who can identify misuses of neural network random seeds and flag them for revision. And of course part of the responsibility falls on NLP authors themselves to make sure they understand the nuances of neural network hyperparameters like random seeds and the ways in which they should and should not be used.

5 Conclusion

I have introduced a simple taxonomy of common uses for neural network random seeds in the NLP literature, describing three appropriate uses (hyperparameter selection, ensemble creation, and sensitivity analysis) and two inappropriate uses (single fixed seed and performance comparison). In an analysis of 85 articles from the ACL Anthology, I have shown that more than half of these recent

NLP articles include inappropriate uses of neural network random seeds. I hope that highlighting this issue can help the NLP community to improve our mentorship and training and move toward more appropriate uses of neural network random seeds in the future.

References

- Jari Björne and Tapio Salakoski. 2018. [Biomedical event extraction using convolutional neural networks and dependency parsing](#). In *Proceedings of the BioNLP 2018 workshop*, pages 98–108, Melbourne, Australia. Association for Computational Linguistics.
- Matt Crane. 2018. [Questionable answers in question answering research: Reproducibility and variability of published results](#). *Transactions of the Association for Computational Linguistics*, 6:241–252.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. [BERT: Pre-training of deep bidirectional transformers for language understanding](#). In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 4171–4186, Minneapolis, Minnesota. Association for Computational Linguistics.
- Jesse Dodge, Gabriel Ilharco, Roy Schwartz, Ali Farhadi, Hannaneh Hajishirzi, and Noah A. Smith. 2020. [Fine-tuning pretrained language models: Weight initializations, data orders, and early stopping](#). *CoRR*, abs/2002.06305.
- Rotem Dror, Gili Baumer, Segev Shlomov, and Roi Reichart. 2018. [The hitchhiker’s guide to testing statistical significance in natural language processing](#). In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 1383–1392, Melbourne, Australia. Association for Computational Linguistics.
- Hang Hua, Xingjian Li, Dejing Dou, Chengzhong Xu, and Jiebo Luo. 2021. [Noise stability regularization for improving BERT fine-tuning](#). In *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 3229–3241, Online. Association for Computational Linguistics.
- Pankaj Kanwar, Reed Wanderman-Milne, and Duncan Riach. 2021. [RFC: Random numbers in TensorFlow 2.0 by wangpengmit - pull request #38 - tensorflow/community](#).
- Adhiguna Kuncoro, Lingpeng Kong, Daniel Fried, Dani Yogatama, Laura Rimell, Chris Dyer, and Phil Blunsom. 2020. [Syntactic structure distillation pre-training for bidirectional encoders](#). *Transactions of the Association for Computational Linguistics*, 8:776–794.

328	Minh Le and Antske Fokkens. 2018. Neural models of		
329	selectional preferences for implicit semantic role la-		
330	beling . In <i>Proceedings of the Eleventh International</i>		
331	<i>Conference on Language Resources and Evaluation</i>		
332	<i>(LREC 2018)</i> , Miyazaki, Japan. European Language		
333	Resources Association (ELRA).		
334	Yan-Jie Lin, Hong-Jie Dai, You-Chen Zhang, Chung-		
335	Yang Wu, Yu-Cheng Chang, Pin-Jou Lu, Chih-Jen		
336	Huang, Yu-Tsang Wang, Hui-Min Hsieh, Kun-San		
337	Chao, Tsang-Wu Liu, I-Shou Chang, Yi-Hsin Con-		
338	nie Yang, Ti-Hao Wang, Ko-Jiunn Liu, Li-Tzong		
339	Chen, and Sheau-Fang Yang. 2020. Cancer registry		
340	information extraction via transfer learning . In <i>Pro-</i>		
341	<i>ceedings of the 3rd Clinical Natural Language Pro-</i>		
342	<i>cessing Workshop</i> , pages 201–208, Online. Associ-		
343	ation for Computational Linguistics.		
344	Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Man-		
345	dar Joshi, Danqi Chen, Omer Levy, Mike Lewis,		
346	Luke Zettlemoyer, and Veselin Stoyanov. 2019.		
347	Roberta: A robustly optimized bert pretraining ap-		
348	proach .		
349	Marius Mosbach, Maksym Andriushchenko, and Diet-		
350	rich Klakow. 2021. On the stability of fine-tuning		
351	BERT: Misconceptions, explanations, and strong		
352	baselines . In <i>International Conference on Learning</i>		
353	<i>Representations</i> .		
354	Garrett Nicolai, Bradley Hauer, Mohammad Motallebi,		
355	Saeed Najafi, and Grzegorz Kondrak. 2017. If you		
356	can’t beat them, join them: the University of Al-		
357	berta system description . In <i>Proceedings of the</i>		
358	<i>CoNLL SIGMORPHON 2017 Shared Task: Univer-</i>		
359	<i>sals Morphological Reinflection</i> , pages 79–84, Van-		
360	couver. Association for Computational Linguistics.		
361	Michael P. Perrone and Leon N. Cooper. 1995. When		
362	networks disagree: Ensemble methods for hybrid		
363	neural networks , pages 342–358. World Scientific.		
364	Nils Reimers and Iryna Gurevych. 2017. Reporting		
365	score distributions makes a difference: Performance		
366	study of LSTM-networks for sequence tagging . In		
367	<i>Proceedings of the 2017 Conference on Empirical</i>		
368	<i>Methods in Natural Language Processing</i> , pages		
369	338–348, Copenhagen, Denmark. Association for		
370	Computational Linguistics.		
371	Julian Risch and Ralf Krestel. 2020. Bagging BERT		
372	models for robust aggression identification . In <i>Pro-</i>		
373	<i>ceedings of the Second Workshop on Trolling, Ag-</i>		
374	<i>gression and Cyberbullying</i> , pages 55–61, Marseille,		
375	France. European Language Resources Association		
376	(ELRA).		
377	Lior Rokach. 2010. Ensemble-based classifiers . <i>Artifi-</i>		
378	<i>cial Intelligence Review</i> , 33(1):1–39.		
379	Shibani Santurkar, Dimitris Tsipras, Andrew Ilyas, and		
380	Aleksander Madry. 2018. How does batch normal-		
381	ization help optimization? In <i>Advances in Neural</i>		
382	<i>Information Processing Systems</i> , volume 31. Curran		
383	Associates, Inc.		
	Two Sigma. 2017. A workaround for non-determinism	384	
	in TensorFlow .	385	
	Alex Wang, Amanpreet Singh, Julian Michael, Fe-	386	
	lix Hill, Omer Levy, and Samuel Bowman. 2018.	387	
	GLUE: A multi-task benchmark and analysis plat-	388	
	form for natural language understanding . In <i>Pro-</i>	389	
	<i>ceedings of the 2018 EMNLP Workshop Black-</i>	390	
	<i>boxNLP: Analyzing and Interpreting Neural Net-</i>	391	
	<i>works for NLP</i> , pages 353–355, Brussels, Belgium.	392	
	Association for Computational Linguistics.	393	
	Sarah Wiegrefe and Yuval Pinter. 2019. Attention is	394	
	not not explanation . In <i>Proceedings of the 2019 Con-</i>	395	
	<i>ference on Empirical Methods in Natural Language</i>	396	
	<i>Processing and the 9th International Joint Confer-</i>	397	
	<i>ence on Natural Language Processing (EMNLP-</i>	398	
	<i>IJCNLP)</i> , pages 11–20, Hong Kong, China. Associ-	399	
	ation for Computational Linguistics.	400	
	Liner Yang and Chencheng Wang. 2019. The BLCU	401	
	system in the BEA 2019 shared task . In <i>Proceedings</i>	402	
	<i>of the Fourteenth Workshop on Innovative Use of</i>	403	
	<i>NLP for Building Educational Applications</i> , pages	404	
	197–206, Florence, Italy. Association for Computa-	405	
	tional Linguistics.	406	