
Transformers as Unrolled Inference in Probabilistic Laplacian Eigenmaps

Anonymous Author(s)

Affiliation

Address

email

Abstract

1 We propose a probabilistic interpretation of transformers as unrolled inference
2 steps, assuming an approximate probabilistic Laplacian Eigenmaps model from
3 the ProbDR framework. Our derivation shows that at initialization, transformers
4 perform “linear” dimensionality reduction. We also show that within the trans-
5 former block, a graph Laplacian term arises from our arguments rather than an
6 attention matrix (which we interpret as an adjacency matrix). We demonstrate
7 that simply subtracting the identity from the attention matrix (and thereby taking a
8 graph diffusion step) improves validation performance on a language model and a
9 simple vision transformer.

10 1 Introduction

11 Transformers, introduced in Vaswani et al. (2017), have been an incredibly successful architecture for
12 deep learning, leading to vastly scaled models used for language as part of Large Language Models
13 (LLMs), such as BERT (Devlin et al., 2019), vision transformers (ViTs) (Dosovitskiy et al., 2021),
14 and foundation models for speech (e.g., wav2vec) (Baevski et al., 2020), as well as in many other
15 application areas.

16 The mathematical basis for their success is an active area of interest. Good generalization cannot be
17 achieved without some assumptions about the underlying data distribution. In this paper, we show
18 that transformers can be seen to perform probabilistic *dimensionality reduction*. Dimensionality
19 reduction enables generalization by imposing a lower dimensional manifold structure on the high
20 dimensional data. Our mathematical approach is heavily inspired by the white-box transformer of Yu
21 et al. (2023), who show that transformers can be viewed as unrolled inference assuming a mixture of
22 Gaussians on latent representations. We provide an alternate interpretation, arguing that each block of
23 a single-head transformer, at initialization, performs gradient descent on a variational lower bound of
24 the probabilistic Laplacian Eigenmaps model of Ravuri et al. (2023). As part of a visual experiment,
25 we show that MNIST flattened images cluster tightly by class when input to a transformer.

26 We use our interpretation of the transformer algorithm to guide us in improving its generalization
27 performance by showing that a modification—simply subtracting an identity matrix from the at-
28 tention matrix (in other words, performing graph diffusion or Laplacian smoothing in the attention
29 step)—follows from our interpretation. We show that this architectural change can be more performant
30 on a language model and vision transformer fit on the tiny Shakespeare and OpenWebText datasets
31 (Karpathy, 2015; Gokaslan et al., 2019), and the downsampled Imagenet datasets (Russakovsky et al.,
32 2015; Chrabaszcz et al., 2017) respectively. This work is a proof of concept on how the insights of
33 Ravuri et al. (2023) can be used to improve engineering tools.

34 Related Work

35 In our work, we interpret the attention matrix as an adjacency matrix of a nearest-neighbor graph
 36 and show that unrolled optimization in a dimensionality reduction model leads to the transformer
 37 architecture.

38 The interpretation of attention matrices as matrices of data-point similarity or relevance has a long
 39 history; Vaswani et al. (2017) and many works since, for instance, Weng (2018); Chefer et al. (2021),
 40 have visualized attention matrices corresponding to text inputs, image patches, etc., for the purposes
 41 of interpretability. Recent work has interpreted the attention matrix as an adjacency matrix and shown
 42 that graph convolutions improve the performance of the architecture (Choi et al., 2024). In our work,
 43 we show that the graph diffusion steps also increase the performance of the architecture. In the realm
 44 of graph convolutional networks, Kipf & Welling (2017) motivate their architecture from a spectral
 45 graph convolutional perspective, and using a slightly different derivation of their updates, we find that
 46 an update also involves a graph Laplacian term of the form $\theta_0 \mathbf{x} + \theta_1 \mathbf{Lx}$. More recently, Joshi (2025)
 47 laid out transformer attention matrices as fully connected graph adjacencies to relate transformers to
 48 graph attention networks of Veličković et al. (2018).

49 Our interpretation of the weight matrices as learning rotation and step size suggests that transformers
 50 learn to learn or learn to optimize quickly (i.e., perform optimization in a latent variable model with
 51 just n_{blocks} steps), which is a well studied field; an overview of the field of learning-to-optimize and
 52 its major ideas is presented in Chen et al. (2021).

53 2 Background

We recap ProbDR’s variational Laplacian Eigenmaps formulation, which forms the basis of our interpretation. Laplacian Eigenmaps is a dimensionality reduction algorithm that reduces the size of a dataset $\mathbf{Y} \in \mathbb{R}^{n \times d}$ to a smaller matrix of representations $\mathbf{X} \in \mathbb{R}^{n \times d_q}$, $d_q \ll d$. The probabilistic Laplacian Eigenmaps model is a probabilistic interpretation of the algorithm (i.e. a model, inference within which leads to the algorithm in question). It can be written as follows, where a Wishart distribution is placed on a precision matrix, of which the graph Laplacian $\mathbf{L}$ is an estimate,

$$d \cdot \mathbf{L}(\mathbf{Y}) \sim \mathcal{W}((\mathbf{X}\mathbf{X}^T + \beta\mathbf{I})^{-1}, d).$$

MAP inference for latent embeddings $\mathbf{X} \in \mathbb{R}^{n, d_q}$ in this model is equivalent to KL minimization over a random variable $\mathbf{\Gamma}$, where the model and variational constraints are written as,

$$\log p(\mathbf{\Gamma}) = \log \mathcal{W}(\mathbf{\Gamma} | (\mathbf{X}\mathbf{X}^T + \beta\mathbf{I})^{-1}, d), \quad \log q(\mathbf{\Gamma}) = \log \mathcal{W}(\mathbf{\Gamma} | \mathbf{L}(\mathbf{Y}), d),$$

where $\mathbf{L}(\mathbf{Y}) \in S_+^n$ is a graph Laplacian¹ matrix encoding a k-nearest neighbour graph, calculated using the data $\mathbf{Y}$. The model graphs are shown in the footnote². It was shown in Ravuri et al. (2023) that the maximum of ELBO, which simplifies as $-\text{KL}(q(\mathbf{\Gamma}) \| p(\mathbf{\Gamma}))$, is attained when the latent embeddings are estimated as follows,

$$\hat{\mathbf{X}} = \mathbf{U}_{d_q} \left(\mathbf{\Lambda}_{d_q}^{-1} - \beta \mathbf{I}_{d_q} \right)^{1/2} \mathbf{R},$$

where $\mathbf{U}_{d_q}$ are the d_q eigenvectors of the graph Laplacian corresponding to the smallest non-zero eigenvalues encoded within the diagonal matrix $\mathbf{\Lambda}$, and where $\mathbf{R} \in O(n)$ is an arbitrary rotation matrix. Further, note that, with an additional constraint, namely $\mathbf{X}^T \mathbf{X} = \mathbf{I}$, the optimal estimate becomes³,

$$\hat{\mathbf{X}} = \mathbf{U}_{d_q} \mathbf{R}.$$

54 In the later case, assuming that the empirical mean of the embeddings is zero, the empirical variance
 55 is equal to $\sum_k \hat{\mathbf{X}}_{kj}^2 / n = 1/n$.

¹We denote the adjacency matrix as $\mathbf{A}$, hence $\mathbf{L} = \mathbf{D} - \mathbf{A}$, with $\mathbf{D}_{ii} = \sum_k \mathbf{A}_{ik}$.

²The model graph can be drawn as: $(\mathbf{x}) \rightarrow (\mathbf{\Gamma})$ and the variational graph as: $(\mathbf{y}) \rightarrow (\mathbf{\Gamma})$.

³This is a consequence of the trace minimisation theorem, as the objective is simply $\text{tr}(\mathbf{L}\mathbf{X}\mathbf{X}^T)$. Any arbitrary rotation still remains a solution as the objective and the constraint are invariant to rotations.

56 The variational interpretation of SimSiam, a Semi-Supervised Learning method

We make a short digression to show how the model graph of ProbDR is not atypical in the representation learning field. Let $\mathbf{Y}_i^a, \mathbf{Y}_i^b, \dots$ be augmentations/views/modalities of a data point. SimSiam, introduced in Chen & He (2020), is a semi-supervised learning method that constructs representations of the data by minimising the negative inner product,

$$\mathcal{L}_i = - \sum_{m_a, m_b} f(h(\mathbf{Y}_i^{m_a}))^\top \textcolor{red}{f}(\mathbf{Y}_i^{m_b}),$$

57 where the element in red is under stop-grad, and with $f(\mathbf{Y}_i^m), f(h(\mathbf{Y}_i^m)) \in \mathcal{S}^{d_q-1}$. Nakamura et al.
58 (2023) show that this loss function has a variational interpretation, where,

$$\begin{aligned} p(\mathbf{X}_i|\mathbf{Y}_i) &\propto \prod_m \text{vMF}(\mathbf{X}_i|f(h(\mathbf{Y}_i^m)), \kappa), & q(\mathbf{X}_i|\mathbf{Y}_i) &\propto \sum_m \delta(\mathbf{X}_i|\textcolor{red}{f}(\mathbf{Y}_i^m)) \\ \Rightarrow \text{KL}(q||p) &= c - \mathbb{E}_{q(\mathbf{X}_i|\mathbf{Y}_i)}(\log p(\mathbf{X}_i|\mathbf{Y}_i)) = - \sum_{m_a, m_b} f(h(\mathbf{Y}_i^{m_a}))^\top \textcolor{red}{f}(\mathbf{Y}_i^{m_b}) = \mathcal{L}_i. \end{aligned}$$

59 Due to the stop-grad applied to the elements of the loss that form the variational constraint, we
60 note that the model graphs are very similar to ProbDR, in that the variational constraint is treated
61 as an observed random variable. We see the variational constraint as approximating a reasonable
62 embedding of the data *at every iteration* of the optimisation process. As an example, if f were
63 initialised as a random projection of the data, it is known that certain properties of the data are
64 retained in the resulting embedding (due to the Johnson–Lindenstrauss lemma). If an optimisation
65 step corresponding to the model preserves/improves these properties (and does not make f degenerate
66 or collapse), we can rely on the variational constraint to always provide an approximate but valid
67 “view” of the data for the model to approach. We apply a similar principle in section 3.

68 Transformers as unrolled optimisation

We now summarise the idea of Yu et al. (2023) on how transformers correspond to unrolled optimisation. Assume a random variable $\mathbf{X} \in \mathbb{R}^{n \times d_q}$, where d_q is the number of latent dimensions and n is the number of i.i.d. data points (of image patches, text tokens, high-dimensional signals, etc.) to which rows of the representations $\mathbf{X}$ correspond. Assuming a latent variable model on $\mathbf{X}$ and a corresponding probabilistic objective $\mathcal{L}$ (a negative log density $-\log p(\mathbf{X})$ or an upper bound on it), gradient descent with m steps of the objective can be unrolled as a sequence of random variables,

$$\mathbf{X}_1 \xrightarrow{T} \mathbf{X}_2 \xrightarrow{T} \dots \xrightarrow{T} \mathbf{X}_s.$$

69 Yu et al. (2023) showed that the gradient descent operation T is very similar to the operations that
70 occur in an (encoder) transformer block, assuming a Gaussian mixture model with a sparse prior
71 on the latent representations $\mathbf{X}$. We note that due to the representations being latent, the model
72 considered in Yu et al. (2023) can also be thought of as a mixture of principal component analysers⁴,
73 therefore suggesting that transformers perform linear (non-kernelized) dimensionality reduction.

74 3 Transformers as ProbDR Inference

In this work, we present an alternative interpretation to that of Yu et al. (2023), that shows that transformers perform gradient descent on a variational objective derived using a variational form of the probabilistic Laplacian Eigenmaps model of Ravuri et al. (2023). We rewrite the random variable corresponding to latents as $\mathbf{Z}$, and treat $\mathbf{X}$ as a parameter that encodes latent positions. We modify the model by adding a prior on the latents,

$$\log p(\mathbf{\Gamma}, \mathbf{Z}) = \log \mathcal{W}(\mathbf{\Gamma} | (\mathbf{Z}\mathbf{Z}^\top + \beta \mathbf{I})^{-1}), d) + \log \mathcal{U}^*(\mathbf{Z}).$$

75 $\mathcal{U}^*$ is a matrix von-Mises-Fisher distribution (a uniform over matrices, with rows that lie on a d_q -
76 dimensional hypersphere), with an additional constraint that for every row $\mathbf{x}$, $\sum_i^{d_q} x_i = 0$ (the rows
77 have zero mean, and hence the coordinates lie on a hyperplane). Projected optimisation with this
78 prior will lead to LayerNorm steps during optimisation.

⁴in a dual sense—acting on the latents and not the components.

We force the random variable $\mathbf{Z}$ to take values $\mathbf{X}$ a.s., and we modify the calculation of the graph Laplacian used in the variational constraint, so that it's a function of the latents $\mathbf{Z}$ and not the data $\mathbf{Y}$,

$$q(\Gamma, \mathbf{Z}) = \mathcal{W}(\Gamma | \tilde{\mathbf{L}}(\mathbf{Z}), d) * \delta(\mathbf{Z} | \mathbf{X}).$$

79 The graph Laplacian is computed as $\tilde{\mathbf{L}} = \mathbf{I} - \tilde{\mathbf{A}}(\mathbf{Z}) = \mathbf{I} - \sigma(\kappa \mathbf{Z} \mathbf{Z}^T - \mathbf{M})$ where σ is the softmax
80 function, applied row-wise (so that the row sums of the input matrix all equal one). $\tilde{\mathbf{A}}$, we argue,
81 is a soft (differentiable) proxy to the true nearest neighbour adjacency matrix, particularly when
82 the latent embeddings $\mathbf{X}$ are initialised with PCA or random projections, as $\mathbf{X} \mathbf{X}^T$ is a minimal-
83 error estimate of the empirical covariance of the data, and the covariance between similar points is
84 expected to be similar in value. This leads to the row-wise softmax being similar and high for similar
85 points, encoding a similarity structure. $\mathbf{M}$ is a mask matrix (for example, if we were to disallow
86 self-adjacency, $\mathbf{M}$ can be set to $\iota \mathbf{I}$, with $\iota \rightarrow \infty$), and κ is a hyperparameter that can be tuned such
87 that the proxy adjacency $\tilde{\mathbf{A}}$ is “close to” a reference nearest neighbour matrix.

88 In a similar fashion to ProbDR, and the variational interpretation to SimSiam, we treat the variational
89 constraint as an observed random variable, and hence do not account for gradient updates to terms
90 leading from the variational constraint. Hence, the KL-div. with stop-grad applied to the variational
91 constraint is,

$$\text{KL}(q(\Gamma, \mathbf{Z}) \| p(\Gamma, \mathbf{Z})) \propto \underbrace{\text{tr}(\tilde{\mathbf{L}}(\mathbf{X} \mathbf{X}^T + \beta \mathbf{I}))}_{\mathcal{L}_{\text{data}}} - \underbrace{\log \det(\mathbf{X} \mathbf{X}^T + \beta \mathbf{I})}_{\mathcal{L}_{\text{reg}}} + c,$$

92 where $\forall i : \mathbf{X}_i \in \mathcal{S}^{d_q-1}$ and $\sum_j \mathbf{X}_{ij} = 0$. Yu et al. (2023) show that a transformer block's
93 sequence of updates follows gradient descent of an objective in steps; given an objective $\mathcal{L}(\mathbf{X}) =$
94 $\mathcal{L}_{\text{data}}(\mathbf{X}) + \mathcal{L}_{\text{reg}}(\mathbf{X})$, they interpret a transformer block calculations as an alternating optimisation
95 process involving the updates,

$$\mathbf{X}' \leftarrow \mathbf{X} - \eta * \frac{d\mathcal{L}_{\text{data}}}{d\mathbf{X}}, \quad \mathbf{X} \leftarrow \mathbf{X}' - \eta * \frac{d\mathcal{L}_{\text{reg}}}{d\mathbf{X}'}.$$

96 In this work, we analyze the transformer at initialization (e.g., with all weights set to diagonal
97 matrices) and consider transformers with single heads, which simplifies the analysis for exposition.
98 We believe that this can be trivially extended by considering a product-of-experts type distribution as
99 part of the variational constraint. Furthermore, for this work, we ignore the ReLU activation that is
100 part of the fully connected segment of the transformer for ease of exposition; however, this can be
101 re-added simply by incorporating a sparsity prior, derived in Yu et al. (2023), as our regularization
102 term is identical to theirs (the sparsity terms notwithstanding).

103 We now show how an (encoder) transformer block's operations arise as optimisation steps of our
104 objective. First, note that, $d\mathcal{L}/d\mathbf{X} = 2\tilde{\mathbf{L}}\mathbf{X} = 2(\tilde{\mathbf{A}} - \mathbf{I})$, and a gradient descent update for optimisation
105 of $\mathcal{L}_{\text{data}}$ follows,

$$\mathbf{X} \leftarrow \mathbf{X} + 2\eta(\sigma(\kappa \mathbf{X} \mathbf{X}^T - \beta \mathbf{I} - \mathbf{M}) - \mathbf{I})\mathbf{X}.$$

106 The element highlighted (which is the degree matrix, in this case, the identity matrix) in red shows
107 the only difference to a standard attention operation (as the attention matrix is the only term that
108 appears in the ordinary architecture). Next, we must take a projection step to ensure that $\forall i : \mathbf{X}_i \in$
109 $\mathcal{S}^{d_q-1}$ and $\sum_j \mathbf{X}_{ij} = 0$, and hence,

$$\mathbf{X} \leftarrow \text{LayerNorm}(\mathbf{X}).$$

110 We now optimise w.r.t. $\mathcal{L}_{\text{reg}}$. Note that this is exactly the same form of regularisation (apart from the
111 sparse prior that gives rise to the ReLU, which is ignored for the sake of exposition, but can trivially
112 be introduced) as the term that appears in Yu et al. (2023). We refer the reader to that work for a
113 careful argument for how this term approximately gives rise to a linear update, but here, we simply
114 approximate $d\mathcal{L}_{\text{reg}}/d\mathbf{X} = 2(\mathbf{X} \mathbf{X}^T + \beta \mathbf{I})^{-1} \mathbf{X} \approx 2/(d_q + \beta) \mathbf{X}$, and our remaining optimisation
115 steps simply involve this update and another projection,

$$\mathbf{X} \leftarrow \mathbf{X} - \frac{2\eta}{\beta + d_q} \mathbf{X}$$

$$\mathbf{X} \leftarrow \text{LayerNorm}(\mathbf{X}),$$

116 which completes the transformer block operations, assuming simple initialisations. Note that a key
117 insight is that the probabilistic interpretation differs from practice in that the former does Laplacian
118 smoothing (**graph diffusion** - i.e. the subtraction of an identity matrix, or a degree matrix, from the
119 attention matrix), whereas the later does not.

120 An interpretation of the weight matrices

121 We posit that an update such as $\mathbf{X} \leftarrow \mathbf{X} + \mathbf{X}\mathbf{W}_{\text{lin}}$ can be interpret as a rotation (which, under the
 122 probabilistic Laplacian Eigenmaps model, the solution is invariant to) and a scaling, which, under
 123 our interpretation, corresponds to a learnt step size $\eta = |\mathbf{W}|^{1/d_q}$; this is a restatement of the belief
 124 that transformers *learn to learn*, in other words, perform optimisation (assuming a dimensionality
 125 reduction or clustering model) with few steps.

126 4 Experiments

127 We provide three main experiments to show validity of the ideas presented thus far. In the first, we
 128 show that a transformer initialised in a simple way performs dimensionality reduction, using flattened
 129 images from the MNIST dataset. In the second, we show that removing an identity matrix from the
 130 attention matrix as suggested by our derivations increases performance on the Shakespeare dataset
 131 and a downsampled (16-by-16) version of Imagenet. In the third, we show that training of GPT-2
 132 converges faster with our modification, than without.

133 4.1 Transformers perform Dimensionality Reduction

134 The details of our dimensionality reduction experiment are as follows. We set up a sequential neural
 135 network, with an initial projection layer with weight $\mathbf{W}_{\text{proj}} \sim \mathcal{MN}(0, \mathbf{I}_d/d, \mathbf{I}_{d_q})$ that is randomly
 136 initialised with Gaussian entries (i.e. a Gaussian random projection). Next, the network consists
 137 of a set of $n_{\text{blocks}} = 8$ encoder transformer blocks. We found that increasing the number of blocks
 138 makes the latents collapse into extremely tight clusters. The LayerNorms have post-normalization
 139 weights associated with them, $\mathbf{W}_{\text{norm}} = \mathbf{I}/\sqrt{n}$, which is because we expect the optimum to be akin
 140 to eigenvectors of a graph Laplacian, which would have variance $1/n$, as explained in the background.
 141 The transformer block weights are $\mathbf{W}_q = \sqrt{\kappa n}\mathbf{I}$, $\mathbf{W}_k = \sqrt{\kappa n/q}\mathbf{I}$ and $\mathbf{W}_v = 2\eta$ corresponding
 142 to the query, key, value weight matrices. The query and key matrices were set up such that the
 143 attention matrix, pre-normalisation, has a diagonal equal to κ . We set $\kappa = 30$, based on the clustering
 144 empirically observed in the resulting graph Laplacian’s eigenvectors. Finally, the feed-forward block
 145 is a single layer with weight $\mathbf{W}_{\text{lin}} = -2\eta$. Note that, based on our derivation (specifically, the scalar
 146 coefficient to the attention matrix), η can be a maximum of 0.5 to avoid magnitudes of the updates
 147 being too large, and so we use the learning rate $\eta = 0.4$. We use the latent dimension $d_q = 128$.
 148 Passing the flattened images through the transformer can be seen to perform clustering, as illustrated
 in fig. 1.



Figure 1: The first two latent dimensions corresponding to flattened MNIST images after a random
 initialisation (i.e. the initial random projection layer that converts pixels to a latent representation)
 (left), and after eight steps through a transformer block (right), showing that transformer blocks
 cluster points in the latent space.

149

150 4.2 Graph Diffusion improves performance

151 In the second experiment, we simply replace the attention matrix $\mathbf{A}$ within a transformer architecture,
 152 found in nanoGPT (Karpathy, 2022) with the negative graph Laplacian $\mathbf{A} - \mathbf{I}$, and run the model
 153 multiple times on the Shakespeare dataset. We also repurpose the code to build a small vision
 154 transformer, and train it naively (i.e. without random augmentations, learning rate schedules, etc.)
 155 on the downsampled Imagenet dataset, wherein all images are 16 by 16 pixels. On this dataset, a

156 benchmark given in Chrabaszcz et al. (2017) achieves 40% validation accuracy, whereas our naïve
 157 ViT achieves around 26%. In both cases however, validation performance improves when we replace
 the attention matrix by the negative graph Laplacian.

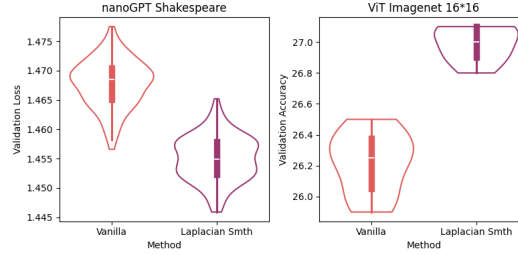


Figure 2: **Left:** validation losses on the Shakespeare dataset and **right:** validation accuracies on a downsampled Imagenet dataset, showing that Laplacian smoothing achieves a better performance in both cases.

158

159 4.3 GPT-2 converges faster with Graph Diffusion

160 In fig. 3, we show the difference in training losses between a run without graph diffusion, and a run
 161 with our modification. We use the same pre-training strategy laid out in Karpathy (2022), however,
 162 we train our GPT-2 model (with about 125M parameters) on a single GH200 GPU (instead of using
 163 torch parallelisation), and we increase batch size and learning rate by a factor of six to make use of
 164 the large memory available.

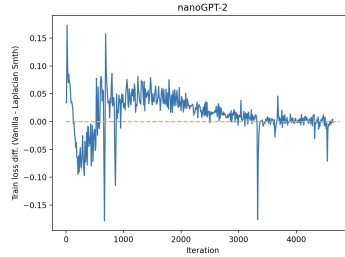


Figure 3: Visualisation of the difference in training losses with and without graph diffusion. A positive difference indicates that the graph diffusion run has achieved a higher performance at any iteration just before convergence. This shows that the graph diffusion run converges slightly faster than the run without any modifications.

165 5 Conclusion

166 We have shown that transformer blocks correspond to unrolled inference assuming a probabilistic
 167 Laplacian Eigenmaps model, and that a simple architectural tweak—using a negative Laplacian $\mathbf{A} - \mathbf{I}$
 168 in place of the attention matrix $\mathbf{A}$ —yields consistent gains in language and vision settings. Future
 169 work will make more careful approximations of the ideas presented, expand on the experimental
 170 validation (current limitations of the work), explore whether non-linear (kernelized) probabilistic
 171 models of dimensionality reduction (from Ravuri & Lawrence (2024)⁵) can increase performance in
 172 models with lower latent dimensionality, and relate transformers to other generalized architectures.
 173 Code used for this paper can be found at (link removed for anonymity) (note that, for the GPT
 174 experiments, this is a very simple modification of Karpathy (2022)).

⁵A simplified version of their objective can be stated as $\text{tr}(\mathbf{L}\mathbf{X}\mathbf{X}^\top) + \sum_{ij} 1/(1 + \|\mathbf{X}_i - \mathbf{X}_j\|^2)$, and it can be shown that an update step with this new regularization term also involves a graph-diffusion-type update.

References

- Alexei Baevski, Henry Zhou, Abdelrahman Mohamed, and Michael Auli. wav2vec 2.0: A framework for self-supervised learning of speech representations, 2020. URL <https://arxiv.org/abs/2006.11477>.
- Hila Chefer, Shir Gur, and Lior Wolf. Transformer interpretability beyond attention visualization, 2021. URL <https://arxiv.org/abs/2012.09838>.
- Tianlong Chen, Xiaohan Chen, Wuyang Chen, Howard Heaton, Jialin Liu, Zhangyang Wang, and Wotao Yin. Learning to optimize: A primer and a benchmark, 2021. URL <https://arxiv.org/abs/2103.12828>.
- Xinlei Chen and Kaiming He. Exploring simple siamese representation learning, 2020. URL <https://arxiv.org/abs/2011.10566>.
- Jeongwhan Choi, Hyowon Wi, Jayoung Kim, Yehjin Shin, Kookjin Lee, Nathaniel Trask, and Noseong Park. Graph convolutions enrich the self-attention in transformers!, 2024. URL <https://arxiv.org/abs/2312.04234>.
- Patryk Chrabaszcz, Ilya Loshchilov, and Frank Hutter. A downsampled variant of imagenet as an alternative to the cifar datasets, 2017. URL <https://arxiv.org/abs/1707.08819>.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. Bert: Pre-training of deep bidirectional transformers for language understanding, 2019. URL <https://arxiv.org/abs/1810.04805>.
- Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, Jakob Uszkoreit, and Neil Houlsby. An image is worth 16x16 words: Transformers for image recognition at scale, 2021. URL <https://arxiv.org/abs/2010.11929>.
- Aaron Gokaslan, Vanya Cohen, Ellie Pavlick, and Stefanie Tellex. Openwebtext corpus. <http://Skylion007.github.io/OpenWebTextCorpus>, 2019.
- Chaitanya K. Joshi. Transformers are graph neural networks, 2025. URL <https://arxiv.org/abs/2506.22084>.
- Andrej Karpathy. char-rnn. <https://github.com/karpathy/char-rnn>, 2015.
- Andrej Karpathy. NanoGPT. <https://github.com/karpathy/nanoGPT>, 2022.
- Thomas N. Kipf and Max Welling. Semi-supervised classification with graph convolutional networks, 2017. URL <https://arxiv.org/abs/1609.02907>.
- Hiroki Nakamura, Masashi Okada, and Tadahiro Taniguchi. Representation uncertainty in self-supervised learning as variational inference. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pp. 16484–16493, 2023.
- Aditya Ravuri and Neil D. Lawrence. Towards one model for classical dimensionality reduction: A probabilistic perspective on umap and t-sne, 2024. URL <https://arxiv.org/abs/2405.17412>.
- Aditya Ravuri, Francisco Vargas, Vidhi Lalchand, and Neil D Lawrence. Dimensionality reduction as probabilistic inference. In *Fifth Symposium on Advances in Approximate Bayesian Inference*, 2023. URL <https://arxiv.org/pdf/2304.07658.pdf>.
- Olga Russakovsky, Jia Deng, Hao Su, Jonathan Krause, Sanjeev Satheesh, Sean Ma, Zhiheng Huang, Andrej Karpathy, Aditya Khosla, Michael Bernstein, Alexander C. Berg, and Li Fei-Fei. Imagenet large scale visual recognition challenge. *International Journal of Computer Vision*, 115(3):211–252, Dec 2015. ISSN 1573-1405. doi: 10.1007/s11263-015-0816-y. URL <https://doi.org/10.1007/s11263-015-0816-y>.

- 220 Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N
221 Gomez, Lukasz Kaiser, and Illia Polosukhin. Attention is all you need. In I. Guyon,
222 U. Von Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett
223 (eds.), *Advances in Neural Information Processing Systems*, volume 30. Curran Associates,
224 Inc., 2017. URL [https://proceedings.neurips.cc/paper_files/paper/2017/file/
225 3f5ee243547dee91fbd053c1c4a845aa-Paper.pdf](https://proceedings.neurips.cc/paper_files/paper/2017/file/3f5ee243547dee91fbd053c1c4a845aa-Paper.pdf).
- 226 Petar Veličković, Guillem Cucurull, Arantxa Casanova, Adriana Romero, Pietro Liò, and Yoshua
227 Bengio. Graph attention networks, 2018. URL <https://arxiv.org/abs/1710.10903>.
- 228 Lilian Weng. Attention? attention! *lilianweng.github.io*, 2018. URL [https://lilianweng.
229 github.io/posts/2018-06-24-attention/](https://lilianweng.github.io/posts/2018-06-24-attention/).
- 230 Yaodong Yu, Sam Buchanan, Druv Pai, Tianzhe Chu, Ziyang Wu, Shengbang Tong, Benjamin D.
231 Haeffele, and Yi Ma. White-box transformers via sparse rate reduction, 2023. URL [https:
232 //arxiv.org/abs/2306.01129](https://arxiv.org/abs/2306.01129).

233 NeurIPS Paper Checklist

234 1. Claims

235 Question: Do the main claims made in the abstract and introduction accurately reflect the
236 paper’s contributions and scope?

237 Answer: [Yes]

238 Justification: The introduction makes careful claims of what is achieved in the paper.

239 2. Limitations

240 Question: Does the paper discuss the limitations of the work performed by the authors?

241 Answer: [Yes]

242 Justification: The limitations are stated briefly in the conclusion.

243 3. Theory assumptions and proofs

244 Question: For each theoretical result, does the paper provide the full set of assumptions and
245 a complete (and correct) proof?

246 Answer: [Yes]

247 Justification: The assumptions are listed clearly, and where results are obtained from external
248 sources, they have been cited.

249 4. Experimental result reproducibility

250 Question: Does the paper fully disclose all the information needed to reproduce the main ex-
251 perimental results of the paper to the extent that it affects the main claims and/or conclusions
252 of the paper (regardless of whether the code and data are provided or not)?

253 Answer: [Yes]

254 Justification: Code provided.

255 5. Open access to data and code

256 Question: Does the paper provide open access to the data and code, with sufficient instruc-
257 tions to faithfully reproduce the main experimental results, as described in supplemental
258 material?

259 Answer: [Yes]

260 Justification: Code provided to run model and visualise data. The data must be obtained
261 independently due to license reasons, but preparation code has also provided (link removed
262 for anonymity).

263 6. Experimental setting/details

264 Question: Does the paper specify all the training and test details (e.g., data splits, hyper-
265 parameters, how they were chosen, type of optimizer, etc.) necessary to understand the
266 results?

267 Answer: [Yes]

268 Justification: Settings needed for understanding the paper have been provided, most other
269 experimental settings are based on Karpathy (2022).

270 7. Experiment statistical significance

271 Question: Does the paper report error bars suitably and correctly defined or other appropriate
272 information about the statistical significance of the experiments?

273 Answer: [Yes]

274 Justification: Error bars provided for smaller experiments, for training GPT-2 it has not been
275 as it’s computationally expensive to run.

276 8. Experiments compute resources

277 Question: For each experiment, does the paper provide sufficient information on the com-
278 puter resources (type of compute workers, memory, time of execution) needed to reproduce
279 the experiments?

280 Answer: [Yes]

281 Justification: The GPU used has been provided, for most other experiments, a much smaller
 282 GPU will suffice.

283 **9. Code of ethics**

284 Question: Does the research conducted in the paper conform, in every respect, with the
 285 NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

286 Answer: [Yes]

287 Justification: Yes.

288 **10. Broader impacts**

289 Question: Does the paper discuss both potential positive societal impacts and negative
 290 societal impacts of the work performed?

291 Answer: [No]

292 Justification: Theoretical research.

293 Guidelines:

294 **11. Safeguards**

295 Question: Does the paper describe safeguards that have been put in place for responsible
 296 release of data or models that have a high risk for misuse (e.g., pretrained language models,
 297 image generators, or scraped datasets)?

298 Answer: [NA]

299 Justification: NA

300 **12. Licenses for existing assets**

301 Question: Are the creators or original owners of assets (e.g., code, data, models), used in
 302 the paper, properly credited and are the license and terms of use explicitly mentioned and
 303 properly respected?

304 Answer: [Yes]

305 Justification: The code that our project is based on is open source and has been cited.

306 **13. New assets**

307 Question: Are new assets introduced in the paper well documented and is the documentation
 308 provided alongside the assets?

309 Answer: [NA]

310 Justification: NA

311 **14. Crowdsourcing and research with human subjects**

312 Question: For crowdsourcing experiments and research with human subjects, does the paper
 313 include the full text of instructions given to participants and screenshots, if applicable, as
 314 well as details about compensation (if any)?

315 Answer: [NA]

316 Justification: NA

317 **15. Institutional review board (IRB) approvals or equivalent for research with human
 318 subjects**

319 Question: Does the paper describe potential risks incurred by study participants, whether
 320 such risks were disclosed to the subjects, and whether Institutional Review Board (IRB)
 321 approvals (or an equivalent approval/review based on the requirements of your country or
 322 institution) were obtained?

323 Answer: [NA]

324 Justification: NA

325 **16. Declaration of LLM usage**

326 Question: Does the paper describe the usage of LLMs if it is an important, original, or
 327 non-standard component of the core methods in this research? Note that if the LLM is used
 328 only for writing, editing, or formatting purposes and does not impact the core methodology,
 329 scientific rigor, or originality of the research, declaration is not required.

330

Answer: [\[Yes\]](#)

331

Justification: Details have been provided on the experimental setup.