



universität
wien

DISSERTATION / DOCTORAL THESIS

Titel der Dissertation /Title of the Doctoral Thesis

„Dimension Reduction with Orthogonal Projections and
Applications in Medical Imaging using Learning Frameworks“

verfasst von / submitted by

Anna Breger, BSc MSc

angestrebter akademischer Grad / in partial fulfilment of the requirements for the degree of
Doktorin der Naturwissenschaften (Dr. rer. nat)

Wien, 2019 / Vienna 2019

Studienkennzahl lt. Studienblatt /
degree programme code as it appears on the student
record sheet:

UA 796 605 405

Dissertationsgebiet lt. Studienblatt /
field of study as it appears on the student record sheet:

Mathematik

Betreut von / Supervisor:

Assoz. Prof. Dr. Martin Ehler, Privatdoz.

Acknowledgement

I want to dedicate this thesis to my beloved grandparents Anton and Elisabeth Fischbach, who have both passed away during the time of my PhD studies. Their intellectualism, kindness and endless love has accompanied and coined my life since my childhood. I am deeply thankful for all their support and inspiration and to have had such wonderful guides in my young life.

I don't want to name here all my loved ones who have supported me in the last years - it wouldn't do justice on how fulfilled and blessed I feel to have so many wonderful people in my life. I am very thankful to have had so much support as well as joyful, deep and inspirational moments and want to earnestly thank everyone who has been part of that!

On the academic side, I want to first sincerely thank my supervisor Martin Ehler, who has given me the freedom to develop and pursue my ideas, but at the same time always was there to answer and discuss my questions, as well as giving valuable advices how to improve my work. I enjoyed and learnt so much in all our critical and demanding discussions. Moreover, I want to thank all my collaborators and colleagues who worked with me and added their expertise, gave new inputs and expanded my horizon.

I am very grateful that I had the chance to meet so many great researchers internationally and especially in the months I spent at the Laboratory of Mathematics in Imaging (LMI) at the Harvard Medical School. Thanks a lot to Carl-Fredrik Westin for providing me the possibility to work there, the stay has been academically and personally very conducive and fruitful. In this breath I also want to thank Katie Mastrogiacono, who helped so much to make all this possible.

Last but not least I want to thank my parents, who have always supported and motivated me. My father as a great scientist has always lived his passion for Maths and Natural Sciences, inspiring me to embrace the challenges of research. My mother as a fantastic musician has always encouraged my passion for music, which provides me an important enhancement to my work as a researcher and at the same time is fostering my creativity.

"Mathematik ist Musik des Geistes, Musik ist Mathematik der Seele"

Daniil Charms (1905 - 1942)

With all that being said, I am very grateful to have had the possibility to work in academic research in the last years and am looking a lot forward to continue and deepen my research activities in the future!

Contents

1	Preamble	1
1.1	Abstract	1
1.2	Introduction	2
1.3	Synopses of the publications	4
1.4	Conclusion	7
2	Publications	13
	(P1) Points on Manifolds with Asymptotically Optimal Covering Radius (<i>Journal of Complexity</i>)	14
	(P2) On Orthogonal Projections for Dimension Reduction and Applications in Augmented Target Loss Functions for Learning Problems (<i>Journal of Mathematical Imaging and Vision</i>)	30
	(P3) Supervised Learning and Dimension Reduction Techniques for Quantification of Retinal Fluid in Optical Coherence Tomography Images (<i>Eye Journal</i>)	51
	<i>Supplementary Material</i>	61
	(P4) An Amplified-Target Loss Approach for Photoreceptor Layer Segmentation in Pathological OCT Scans (<i>Lecture Notes in Computer Science, Springer</i>)	64
	(P5) On the Reconstruction Accuracy of Multi-Coil MRI with Orthogonal Projections (<i>Preprint</i>)	74
A	Appendix	89
A.1	Zusammenfassung	89
A.2	Curriculum Vitae	90

1 Preamble

1.1 Abstract

In the age of digitalized data, dimension reduction is a necessary tool to compress and combine high-dimensional data. The lower dimensional representations shall enable the feasibility of computations, but at the same time preserve essential information for subsequent tasks. In this thesis I derive new mathematical results regarding dimension reduction with orthogonal projections and connect them to problems arising in clinical medical routine, enabling new analyses and improving performances in (deep) learning tasks.

Orthogonal projections are important linear methods in that framework, since the dimension reduction itself is just based on a matrix multiplication, allowing fast computation on large data sets. We derive new results on point sequences that sample the space of orthogonal projections, called the Grassmannian. In particular, we prove that a specific class of sequences covers the space in an asymptotically optimal manner. Our numerical experiments illustrate the theoretical results, validating further analyses and experiments in applications with such point sequences.

To measure the adequacy of a dimension reduction method, the preservation of diverse aspects of information must be considered. The orthogonal projection arising from the well known principal component analysis (PCA) maximizes the total variance within the projected data. Especially tasks that work with noisy data benefit from this method, since the influence of noise can be reduced by cutting off the eigendirections that correspond to the smaller principal components. On the other hand, random projections ensure with high probability the preservation of all pairwise distances within a data set. This property is especially important for tasks that rely on small changes within the data. Whether one or the other kind of information is of main importance depends on further objectives and the condition of the given data. We provide a mathematical proof to show that these two objectives strongly correlate and usually cannot be achieved at the same time. Numerical experiments are designed to illustrate these results. Moreover, we determine specific projections on input data in a standard classification problem. The highest classification accuracy was achieved by projections that balance the two objectives, preservation of variance and pairwise distances.

Dimension reduction and orthogonal projections are useful in plenty diverse imaging problems. In the applied part of my PhD project I work with clinical optical coherence tomography (OCT) and magnetic resonance imaging (MRI) data. This research emerged from collaborations with the Vienna Reading Center (VRC), Department of Ophthalmology, Medical University of Vienna as well as the Laboratory of Mathematics in Imaging (LMI), Brigham and Women's Hospital, Harvard Medical School, Boston, MA. Automated image segmentation is especially important and necessary for disease identification, since daily routine of doctors

usually does not allow enough time to look through all scans of patients. Based on annotated OCT data we develop image segmentation methods using (deep) learning algorithms that include dimension reduction and orthogonal projections. Our methods improve the accuracy of baseline methods and could serve as future guidelines for doctors. Moreover, in an analysis on MRI reconstruction accuracy, we use orthogonal projections to linearly compress higher dimensional signal arrays and conclude the beneficial impact of PCA.

1.2 Introduction

The technological development of the last decade has enabled fast detailed acquirement and storage of high-dimensional digitalized data. Massive amounts of data are collected daily in numerous different areas of modern life and need to be analyzed. This created not only the occupational field of data scientists, but also revolutionized and expanded traditional research fields immensely, such as mathematical data analysis. The fast technological evolutions open constantly lots of new possibilities and challenges in different areas that work with digitalized data. One of the areas that has highly benefited from the technological developments is medical imaging, including the acquirement of patient data in hospitals itself as well as subsequent analyses.

Medical image analysis is of main importance for several clinical applications, such as disease recognition, identification and quantification of abnormalities, as well as tracking of disease progression. Automated tools have the potential to help doctors make much faster and more precise diagnoses and decisions. It is often very challenging and time consuming to manually annotate and evaluate data, even for medical experts. Moreover, studies have shown that inconsistencies between trained graders are normal to a certain percentage (see e.g. [36]). Therefore, sufficient automatization of annotation tasks, such as classification and quantification, is very important to gain more consistency and enable the evaluation of a large amount of patient data. Arising problems with clinical data are usually highly complex and applying standard imaging methods are often not very successful. Convolutional neural networks (CNN) have become the standard choice in many medical image contexts, such as segmentation tasks [32]. The training of a neural network is based on a loss function that is optimized. Besides the choices of hyperparameters (number of hidden units, dropout rate, regularization, ...), an adequate adaption of the loss function itself can improve the performance in complicated medical tasks, see e.g. [30, 19].

In this thesis I aim to establish theoretical results on orthogonal projections for dimension reduction that shall be illustrated by numerical experiments. Moreover, the mathematical insights shall be applied in clinical settings with medical image data, enabling new analyses and improving automated image segmentation.

Computational cost of many data analysis methods, such as various machine learning tools, depend on the dimension of the data. Therefore lower dimensional representations are needed that enable the feasibility of computations, but at the same time preserve or even augment essential information for subsequent tasks. Linear dimension reduction is a powerful tool for the preprocessing of high-dimensional data to compress and combine important data features

[12]. In contrast to nonlinear approaches (such as [34, 33, 16, 13]), the use of orthogonal projections is computationally cheap, since it corresponds to a simple matrix multiplication.

We aim to identify orthogonal projections that preserve essential information and complexity within a more compact representation. The preservation of pairwise distances is one kind of information property that can be measured within the projected lower-dimensional data set. The Johnson-Lindenstrauss Lemma ([23]) yields that there exists such an orthogonal projector that preserves all pairwise distances up to some small error margin. Its proof reveals that the random choice of projectors works with high probability, cf. [14]. Methods with random projections have been advantageously used to reduce the dimension on image data, see e.g. [18, 2]. On the other hand, variants of the principal component analysis (PCA) are widespread and have been successfully applied in many tasks, including numerous medical imaging problems [28]. PCA directly gains another kind of information preservation, since it yields an orthogonal projection that maximizes the variance within the projected data [26]. It also reveals a denoising effect, which can act beneficial on noisy image data, see e.g. [15, 27]. Regarding computational costs, the use of random projections is especially favorable for large, high-dimensional data ([35]), since the construction is independent from the actual data. The computational complexity is just $O(dkm)$, e.g. using the construction in [1], with $d, k \in \mathbb{N}$ being the original and lower dimensions and $m \in \mathbb{N}$ the number of samples. In contrast, PCA needs much more operations, i.e. $O(d^2m) + O(d^3)$, cf. [20]. However, it is not obvious what kind of information preservation is generally most beneficial. Both objectives, preservation of either pairwise distances or variance, appear useful and as stated above, act beneficially in different settings. The question arises how the two preservation properties interact and in what sense they are competing.

In order to study orthogonal projections for dimension reduction in numerical experiments, we need a finite point sequence (i.e. finite covering) that well represents the underlying space. For fixed d and k the space of orthogonal projections is called the Grassmannian. One way to measure the efficiency of a covering is by analyzing its covering radius in relation to its cardinality. The covering radius describes the biggest hole in a given point set and therefore we aim for its minimization. To enable further usage in experiments, we want to determine theoretically optimal covering sequences that also can be constructed numerically.

Building upon our theoretical results, we aim to connect to medical imaging tasks via derivation of new analysis schemes and adaption of automatization methods. In collaboration with the Vienna Reading Center (VRC) at the Department of Ophthalmology, Medical University of Vienna, we work with clinical optical coherence tomography (OCT) images of the human retina. In addition to the patient scans they provided manual annotations performed by trained graders, serving as ground truth and enabling supervised learning. Important tasks in clinical OCT images are photoreceptor segmentation and the quantification of retinal fluid. The photoreceptors are a single cell layer that is capable of visual phototransduction. Disruptions within this layer highly correlate with focal vision impairment and can be caused by retinal fluid. Thereby, OCT provides good visualization, enabling identification and registration of fluid appearance between patient visits, which is important for therapy. Moreover, direct photoreceptor segmentation and in particular the identification of disruptions can highlight abnormalities. Segmentation models with standard convolutional neural networks often do not yield sufficient results on pathological patient data. We aim to improve the automatization

with dimension reduction tools, orthogonal projections, and the incorporation of important target data characteristics.

In collaboration with the Laboratory of Mathematics in Imaging (LMI) at the Harvard Medical School, we work with magnetic resonance images (MRI). Signal acquisition from multiple coils/receivers is nowadays commonplace in clinical MRI routine, asking for array combination in the image reconstruction pipeline. We aim to analyze this step in terms of standard methods and linear compression. Random orthogonal projections enable a correlation analysis of important performance measures, investigating the impact of linear compression on the final image quality.

The described mathematical and medical objectives lead to the following key questions that I studied within my PhD project:

- (A) Can we numerically construct coverings in the Grassmannian that are optimal in some sense?
 - (A1) Derive theoretical tools suitable for numerical constructions.
 - (A2) Develop numerical experiments supporting the theoretical results.
- (B) Can we use specific orthogonal projections for dimension reduction in a beneficial way?
 - (B1) Identify important information preservation properties in dimension reduction and derive theoretical results that connect these properties.
 - (B2) Design associated numerical experiments with Grassmannian coverings and classical learning tools.
- (C) Can we use dimension reduction and data compression with orthogonal projections in medical imaging tasks?
 - (C1) Use dimension reduction and learning tools for automated retinal fluid quantification.
 - (C2) Improve photoreceptor segmentation in retinal images with CNNs through adaption of the loss function including orthogonal projections.
 - (C3) Apply orthogonal projections to coil combination in multidimensional reconstruction of magnetic resonance brain imaging data.

In the following subsection I will briefly describe the content of the individual papers that are contained in this thesis (P1 - P5), followed by the conclusion section that discusses how the key questions (A), (B) and (C) have been answered within the papers.

1.3 Synopses of the publications

P1 Points on manifolds with asymptotically optimal covering radius [8]

Journal of Complexity, Elsevier

The covering radius of a finite set of points can quantify their covering properties: the

smaller the radius the better is the underlying space covered. The worst case quadrature error in Sobolev spaces bounds the covering radius for points on the Euclidean sphere. Here, we extend the results to points on compact smooth Riemannian manifolds. Moreover, the terminology of asymptotically optimal coverings is defined. We show then, that specific points on compact smooth Riemannian manifolds achieve this asymptotically optimal covering radius. The Grassmannian serves us as an important example for a compact smooth Riemannian manifold. We aim to numerically construct such asymptotically optimal covering sequences to enable numerical experiments that demonstrate our theoretical findings in the Grassmannian. However, the numerical computation of the covering radius for a given point set is generally infeasible, since it asks for combinatorial optimization over the whole Grassmannian. We overcome this issue by using a very large amount of random points to represent the whole space. To show that this choice is sufficiently accurate, we combine the known asymptotics for the covering radius of random points and the canonical volumes induced by the Euclidean metric to explicitly compute their covering radius. With that we can choose the amount of random points such that the random covering reliably represents the whole Grassmannian. Finally, together with the asymptotically optimal covering sequences, we can provide numerical experiments that illustrate our theoretical results.

P2 On orthogonal projections for dimension reduction and applications in augmented target loss functions for learning problems [10]

Journal of Mathematical Imaging and Vision (JMIV), Springer

In this paper we investigate the relation between two important information preservation properties in dimension reduction: preservation of pairwise distances and preservation of total variance (which is equivalent to total variance maximization of the projected data). The Johnson-Lindenstrauss Lemma tells that with high probability all pairwise distances within a randomly projected data set will be preserved. Hence, random projections yield this kind of information preservation with high probability. On the other hand, one of the most frequently used linear dimension reduction methods, principal component analysis (PCA), maximizes the total variance within the projected data set. Now the question occurs if it is possible to gain both preservation properties. In the main theorem we determine the correlation between those two properties and the asymptotics imply that usually it is not possible to achieve both preservation properties at the same time. In a toy classification experiment we reduce the dimension of the input data with different projections and compare the results. The best results were achieved with a projection that balances both properties. This suggests that the choice of information preservation should be included in the decision of the dimension reduction method.

In the second part of the paper we introduce the framework of augmented target (AT) loss functions for learning problems. Conventional loss functions in supervised learning with neural networks compare the outputs in each iteration with some given ground truth. Often, there are known characteristics about the target space that are fairly unused, since their representation in the target data is not predominant. Therefore other characteristics are favored in the optimization process. The AT loss functions incorporate prior knowledge via transformations of the target space. Comparing the outputs and

targets in this transformed spaces adds new information to the optimization problem that can enhance the performance of the neural network training. Which transformations are meaningful depend on the particular tasks and orthogonal projections can help to weight them. Here we show two experiments: an image segmentation problem with medical data and a musical instrument classification task. In both cases the adjusted AT loss function can enhance the performance.

P3 Supervised learning and dimension reduction techniques for quantification of retinal fluid in optical coherence tomography images [6]

Eye (Journal), Springer Nature

Dimension reduction, traditional image segmentation, and supervised learning methods are used in this paper to create a pipeline for retinal fluid quantification in clinical optical coherence tomography (OCT) images. Serving as ground truth, the image volumes are annotated manually in each cross-section by human expert graders of the Vienna Reading Center (VRC). Automated fluid quantification can help doctors to accelerate evaluations in every day use, where the identification of fluid and the tracking of progression between two visits of a patient is of high importance.

Our fully automated pipeline is based on good pre-segmentation results from applying an energy optimization model to vector-valued feature images. These feature channels were chosen as a texture enhancing and a smoothed image, resulting from another variational energy minimization. In the last step of the pipeline, a threshold is learned based on the ground truth and identifies the region of retinal fluid. The beneficial interplay of the methods yields good results, visually and in numbers. It appears well-suited for integration into actual manufacturer's devices, facilitating image segmentation in daily clinical routine.

P4 An amplified-target loss approach for photoreceptor layer segmentation in pathological OCT scans [9]

Ophthalmic Medical Image Analysis (OMIA), Lecture Notes in Computer Science, Springer

In this paper we apply an adjusted augmented target loss function, that we introduced in a general framework in P2, to photoreceptor layer segmentation with convolutional neural networks (CNN).

Supervised deep learning models trained with standard loss functions are usually able to characterize only the most common disease appearance from a training set. This often results in suboptimal performance and poor generalization when dealing with unseen lesions. Our newly designed loss function explicitly penalizes errors within the central area of the input images, based on the observation that most of the challenging disease appearance is located in this area. The increased performance in the clinical test set suggests that our adapted loss function improves the identification of photoreceptors in highly pathological scenarios.

P5 On the reconstruction accuracy of multi-coil MRI with orthogonal projections [11]

Preprint, arXiv

We study the reconstruction accuracy regarding linear compression of magnetic resonance images (MRI) from multiple coils/receivers. Signal acquisition from multiple receivers is nowadays a standard setting in clinical routine, due to better and faster imaging. The whole reconstruction pipeline from the signal of the scanner to the final images includes a sampling step on the measured data. Some standard reconstruction methods yield multidimensional signal arrays that have to be combined in one final image. Optimal array combination relies on information from the scanning process, which is not always available. To overcome the dependence on prior knowledge, often the root sum of squares (rSOS) is used to combine the image channels. Although this straightforward method yields relatively good results for low noise levels, data that contains artifacts and a high amount of noise is usually poorly reconstructed.

The covering property of random orthogonal projections enables us to numerically analyze the reconstruction quality when including linear compressions in the coil combination step. To evaluate the reconstruction of the adjusted coil combination, we consider diverse measures of accuracy: L_2 reconstruction error, variance, signal-to-noise-ratio (SNR) and visual evaluation. In numerical experiments we study the coil combination with two data sets, a simulated with varying noise levels and an in-vivo data set. The resulting reconstructions show strong correlation between variance and L_2 reconstruction error, but do not necessarily relate to the visual assessment. The SNR as a measure of quality relates more accurately to the visual evaluation. Benefiting from the denoising effect, PCA yields the highest SNR in all experiments. These results suggest to include linear compression of image space data with PCA.

1.4 Conclusion

The key question (A) has been fully answered in P1 ([8]). (A1) was examined by extending results of asymptotically optimal coverings for quasi-Monte Carlo points on the Euclidean sphere to points on compact smooth Riemannian manifolds. On the Euclidean sphere it was shown in [4] that the worst case error of integration bounds the covering radius. Moreover, quasi-Monte Carlo integration points on the sphere have been connected to the worst case error of integration, cf. [5]. We combine these results with general findings in [3] on the worst case error on compact smooth Riemannian manifolds. This leads to our theorem, stating that quasi-Monte Carlo integration points on compact smooth Riemannian manifolds provide asymptotically optimal covering radii. In particular this holds for the space of our further interest, the Grassmannian. Following the numerical construction through minimization in [17, 7] provides such asymptotically optimal sequences of orthogonal projectors. The worst case error of integration can be calculated analytically on specific reproducing kernel Hilbert spaces, cf. [7, 21, 29]. This enables the computation of the point sequences for further numerical experiments. Based on properties of random projections, we newly design experiments that illustrate and approve our theoretical findings, solving (A2). The mathematical and numerical results on the covering radius indicate that the asymptotically optimal coverings of the Grassmannian can yield finite sequences of orthogonal projections that well represent the whole underlying space. This serves us as a basis for the subsequent research on dimension

reduction with orthogonal projections, including applications to medical data.

In P2 ([10]) we studied important information preservation properties in dimension reduction with orthogonal projections. Popular choices thereby are variants of PCA, see e.g. [28, 15] and random projections, see e.g. [2, 35], with their underlying characteristics to closely preserve total variance and pairwise distances respectively, cf. [26, 23]. We introduce quantities that describe these two information preservation properties. Known characteristics of orthogonal projections enable the computation of statistical relations between these quantities and yield correlation bounds and asymptotics. The results indicate that both information properties usually cannot be achieved at the same time and one has to favor or balance the objectives. With that (B1) is answered, focusing on these two important examples of preservation properties. Next, we provide numerical studies that compare different lower dimensional input representations in classification experiments on a popular publicly available data set. In particular, relying on the optimal covering property observed in P1, we use a finite asymptotically optimal covering sequence to choose specific orthogonal projections. The classification results, after applying these projections to the input data, suggest projections that balance the two preservation objectives. With that (B2) is addressed, solving the problems occurring in (B). Generally, the theoretical and numerical results indicate that sufficient linear compression needs to take care of an appropriate choice of information preservation properties. Since not all kinds of preservation can usually be achieved at the same time, balancing projections can give a beneficial alternative in scenarios where different properties are needed and conventional methods would lose too much information.

The applications in medical imaging with clinical data (C) ask for advanced automatization and analysis methods. This involves improvement of complicated retinal image segmentation tasks with OCT data. Supervised deep learning techniques are state of the art for medical image segmentation [25], particularly with fully convolutional neural network architectures such as the U-Net [32]. A loss function has to be optimized to train the neural network and affects the outcome directly. Adapted loss functions that include target features have been successfully applied to different imaging tasks, such as super-resolution [24], transfer learning [22] and image segmentation [30, 19]. We introduced in P2 a framework of augmented target loss functions that enable the incorporation of prior knowledge via transformations. Orthogonal projections can weight these transformations, that serve as additional target information features. For photoreceptor segmentation with a convolutional neural network, we formulate such a loss function that emphasizes edge detection by projected features. This new loss function improves the accuracy of the segmentation in pathological image patient data. This solves (C2) together with P4 ([9]), where we designed an augmented target loss function that amplifies important retinal regions and applied it to photoreceptor segmentation on OCT volumes with three common macular diseases. The new loss function improved especially the segmentation of those parts in the image volumes where the most disease related abnormalities appear. Furthermore, we used dimension reduction tools to quantify retinal fluid successfully, see P5 ([6]). Our fully automatized segmentation pipeline seems well-suited for integration in daily clinical routine to support the evaluation of patient data, answering (C1).

Orthogonal projections for linear data combination were studied in P5 ([11]) in the context of magnetic resonance imaging (MRI), addressing (C3). Following the ideas of good covering in P1 and the correlation analysis of important compression properties in P2, we use random

orthogonal projections as a basis to analyze the reconstruction accuracy. We include linear compression in coil combination for multi-coil MRI data, where root sum of squares (rSOS) is the standard method when no detailed knowledge about the coils is available ([31, 37]). This additional weighting yields new representations of the coil images, following our idea of feature combination in P2. We observe the results with diverse image quality measures, including visual evaluation. Our numerical analyses on simulated and in-vivo data with different noise levels affirm to incorporate PCA as linear compression in the image domain.

Summing up, I connected our mathematical results on orthogonal projections to applications in medical imaging, enabling improvements and better understanding in complicated clinical tasks. The mathematical basis provides new possibilities and views, as well as challenges in direct applications. Future research shall include the deterministic finding of projections that balance different information preservation properties. Beyond the studied information properties, we aim to investigate in cooperation with clinicians more closely how diverse quality measures can describe medical image data. Amongst others, this shall serve as a basis for unbiased and accurate evaluation of dimension reduction methods in clinical applications.

Our framework of augmented target loss functions gives the possibility to incorporate prior knowledge in learning tasks. Problems arising in medical imaging are usually highly complex and data representations that facilitate important characteristics can be incorporated in our loss function's framework. Following the great potential and the success in photoreceptor segmentation, we plan to apply augmented target loss functions to other medical imaging tasks, such as retinal fluid quantification in OCT or brain tissue segmentation in MRI data. This shall be investigated in the future by continuing and extending current collaborations.

Bibliography

- [1] D. Achlioptas. Database-friendly random projections: Johnson-Lindenstrauss with binary coins. *Journal of Computer and System Sciences*, 66(4):671–687, 2003.
- [2] E. Bingham and H. Mannila. Random projection in dimensionality reduction: Applications to image and text data. In *Proceedings of the Seventh ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, KDD '01, pages 245–250, New York, NY, USA, 2001. ACM.
- [3] L. Brandolini, C. Choirat, L. Colzani, G. Gigante, R. Seri, and G. Travaglini. Quadrature rules and distribution of points on manifolds. *Annali della Scuola Normale Superiore di Pisa - Classe di Scienze*, XIII(4):889–923, 2014.
- [4] J. S. Brauchart, J. Dick, E. B. Saff, I. H. Sloan, Y. G. Wang, and R. S. Womersley. Covering of spheres by spherical caps and worst-case error for equal weight cubature in Sobolev spaces. *Journal of Mathematical Analysis and Applications*, 431(2):782–811, 2015.
- [5] J. S. Brauchart, E. Saff, I. H. Sloan, and R. Womersley. QMC designs: Optimal order quasi Monte Carlo integration schemes on the sphere. *Mathematics of Computation*, 83:2821–2851, 2014.
- [6] A. Breger, M. Ehler, H. Bogunovic, S. M. Waldstein, A.-M. Philip, U. Schmidt-Erfurth, and B. S. Gerendas. Supervised learning and dimension reduction techniques for quantification of retinal fluid in optical coherence tomography images. *Eye*, 31(8):1212–1220, 2017.
- [7] A. Breger, M. Ehler, and M. Gräf. *Quasi Monte Carlo integration and kernel-based function approximation on Grassmannians*. Applied and Numerical Harmonic Analysis series (ANHA). Birkhauser/Springer, 2017.
- [8] A. Breger, M. Ehler, and M. Gräf. Points on manifolds with asymptotically optimal covering radius. *Journal of Complexity*, 48:1–14, 2018.
- [9] A. Breger, J. I. Orlando, H. Bogunović, S. Riedl, B. Gerendas, M. Ehler, and U. Schmidt-Erfurth. An amplified-target loss approach for photoreceptor layer segmentation in pathological oct scans. In *Lecture Notes in Computer Science*, volume 11855, pages 26 – 34. Ophthalmic Medical Image Analysis (OMIA), Springer, 2019.
- [10] A. Breger, J. I. Orlando, P. Harár, M. Doerfler, S. Klimscha, C. Grechenig, B. Gerendas, U. Schmidt-Erfurth, and M. Ehler. On orthogonal projections for dimension reduction and applications in augmented target loss functions for learning problems. *Journal of Mathematical Imaging and Vision (JMIV)*, Springer, 2019.
- [11] A. Breger, G. Ramos Llorden, G. Vegas Sanchez Ferrero, W. S. Hoge, M. Ehler, and C.-F. Westin. On the reconstruction accuracy of multi-coil mri with orthogonal projections. arXiv:1910.13422, 2019.
- [12] J. P. Cunningham and Z. Ghahramani. Linear dimensionality reduction: Survey, insights, and generalizations. *Journal of Machine Learning Research*, 16:2859–2900, 2015.
- [13] W. Czaja and A. Halevy. On convergence of schroedinger eigenmaps. *preprint*, 2011.
- [14] S. Dasgupta and A. Gupta. An elementary proof of a theorem of Johnson and Lindenstrauss. *Random Structures & Algorithms*, 22(1):60–65, 2003.
- [15] M. D. Does, J. L. Olesen, K. D. Harkins, T. Serradas-Duarte, D. F. Gochberg, S. N. Jespersen, and N. Shemesh. Evaluation of principal component analysis image denoising on multi-exponential mri relaxometry. *Magnetic Resonance in Medicine*, 81(6):3503–3514, 2019.

- [16] D. L. Donoho and C. Grimes. Hessian eigenmaps: new locally linear embedding techniques for high-dimensional data. *Proceedings of the National Academy of Sciences*, 100:5591–5596, 2003.
- [17] M. Ehler and M. Gräf. Reproducing kernels for the irreducible components of polynomial spaces on unions of Grassmannians. *Constructive Approximation (Springer)*, 49(29-58), 2018.
- [18] X. Z. Fern and C. E. Brodley. Random projection for high dimensional data clustering: A cluster ensemble approach. In *Proceedings of the Twentieth International Conference on International Conference on Machine Learning*, ICML’03, pages 186–193. AAAI Press, 2003.
- [19] P.-A. Ganaye, M. Sdika, and H. Benoit-Cattin. Semi-supervised learning for segmentation under semantic constraint. pages 595–602, Spain, 2018. MICCAI.
- [20] G. H. Golub and C. F. V. Loan. *Matrix Computations*. Johns Hopkins Studies in the Mathematical Sciences. The Johns Hopkins University Press, 1996.
- [21] M. Gräf. *Efficient Algorithms for the Computation of Optimal Quadrature Points on Riemannian Manifolds*. Universitätsverlag Chemnitz, 2013.
- [22] J. Johnson, A. Alahi, and L. Fei-Fei. Perceptual losses for real-time style transfer and super-resolution. In *ECCV*, 2016.
- [23] W. Johnson and J. Lindenstrauss. Extensions of lipschitz maps into a hilbert space. *Contemporary Mathematics*, 26:189–206, 01 1984.
- [24] C. Ledig, L. Theis, F. Huszar, J. Caballero, A. P. Aitken, A. Tejani, J. Totz, Z. Wang, and W. Shi. Photo-realistic single image super-resolution using a generative adversarial network. *2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 105–114, 2017.
- [25] G. Litjens, T. Kooi, B. E. Bejnordi, A. A. A. Setio, F. Ciompi, M. Ghafoorian, J. A. W. M. van der Laak, B. van Ginneken, and C. I. Sánchez. A survey on deep learning in medical image analysis. *Medical image analysis*, 42:60–88, 2017.
- [26] B. Moore. Principal component analysis in linear systems: Controllability, observability, and model reduction. *IEEE Transactions on Automatic Control*, 26, 1981.
- [27] Y. Murali, M. Babu, D. Subramanyam, and D. Prasad. Pca based image denoising. *Signal And Image Processing*, 3, 04 2012.
- [28] D. Nandi, A. S. Ashour, S. Samanta, S. Chakraborty, M. A.-M. M. Salem, and N. Dey. Principal component analysis in medical image processing: A study. *International journal of Image Mining*, 1:45–, 01 2015.
- [29] E. Novak and H. Wozniakowski. *Tractability of Multivariate Problems. Volume II*, volume 12 of *EMS Tracts in Mathematics*. EMS Publishing House, Zürich, 2010.
- [30] O. Oktay, E. Ferrante, K. Kamnitsas, M. Heinrich, W. Bai, J. Caballero, R. Guerrero, S. A Cook, A. de Marvao, T. Dawes, D. O’Regan, B. Kainz, B. Glocker, and D. Rueckert. Anatomically constrained neural networks (acnn): Application to cardiac image enhancement and segmentation. *IEEE Transactions on Medical Imaging*, PP, 05 2017.
- [31] P. B. Roemer, W. A. Edelstein, C. E. Hayes, S. P. Souza, and O. M. Mueller. The NMR phased array. *Magnetic Resonance in Medicine*, 16(2):192–225, 1990.
- [32] O. Ronneberger, P. Fischer, and T. Brox. U-Net: Convolutional Networks for Biomedical Image Segmentation. In N. Navab, J. Hornegger, W. M. Wells, and A. F. Frangi, editors, *MICCAI 2015*, pages 234–241. Springer, 2015.
- [33] S. T. Roweis and L. K. Saul. Nonlinear dimensionality reduction by locally linear embedding. *Science*, 290(5500):2323–2326, 2000.
- [34] J. B. Tenenbaum, V. de Silva, and J. C. Langford. A global geometric framework for nonlinear dimensionality reduction. *Science*, 290(5500):2319–23, Dec 2000.

- [35] G.-A. Thanei, C. Heinze, and N. Meinshausen. Random projections for large-scale regression. In *Big and Complex Data Analysis: Methodologies and Applications*, pages 51–68, Cham, 2017. Springer International Publishing.
- [36] E. S. Varnousfaderani, J. Wu, W.-D. Vogl, A.-M. Philip, A. Montuoro, R. Leitner, C. Simader, S. M. Waldstein, B. S. Gerendas, and U. Schmidt-Erfurth. A novel benchmark model for intelligent annotation of spectral-domain optical coherence tomography scans using the example of cyst annotation. *Computer Methods and Programs in Biomedicine*, 130:93–105, 2016.
- [37] D. O. Walsh, A. F. Gmitro, and M. W. Marcellin. Adaptive reconstruction of phased array mr imagery. *Magnetic Resonance in Medicine*, 43(5):682–690, 2000.

2 Publications

(P1)

Points on Manifolds with Asymptotically Optimal Covering Radius

Journal: Journal of Complexity, Elsevier

Authors: A. Breger, M. Ehler and M.Gräf

Published: 2018

Links: www.sciencedirect.com/science/article/pii/S0885064X18300256

arXiv: <https://arxiv.org/abs/1607.06899>

DOI: 10.1016

Contribution: Theory development; Design and performance of the numerical experiments; Writing.



Contents lists available at ScienceDirect

Journal of Complexity

journal homepage: www.elsevier.com/locate/jco

Points on manifolds with asymptotically optimal covering radius[☆]

Anna Breger^{a,*}, Martin Ehler^{a,*}, Manuel Gräf^b

^a University of Vienna, Department of Mathematics, Oskar-Morgenstern-Platz 1 A-1090 Vienna, Austria

^b Austrian Academy of Sciences, Acoustics Research Institute, Vienna, Austria



ARTICLE INFO

Article history:

Received 22 July 2016

Accepted 4 March 2018

Available online 30 March 2018

Keywords:

Covering radius

Worst case integration error

Riemannian manifold

Cubature points

Quasi-Monte Carlo integration

ABSTRACT

Given a finite set of points on the Euclidean sphere, the worst case quadrature error for functions in Sobolev spaces has recently been shown to provide upper bounds on the covering radius of the point set. Moreover, sequences of Quasi-Monte Carlo (QMC) designs for Sobolev spaces on the sphere achieve asymptotically optimal covering radii. Here, we extend these results from sequences of QMC designs on the sphere to sequences of weighted QMC designs on compact smooth Riemannian manifolds. We also provide numerical experiments illustrating our findings for the Grassmannian manifold.

© 2018 Elsevier Inc. All rights reserved.

1. Introduction

Many discretization schemes in numerical analysis are based on finite samples that cover the underlying space sufficiently well, i.e., the sampling points have small covering radius. One way to measure the efficiency of a family of coverings is to determine the asymptotic behavior, specifically, the rate of convergence of the covering radius as the cardinality of the coverings tends to infinity.

According to [8] the covering radius of an n -point set on the unit d -sphere is bounded above by a power of the associated worst-case error for equal weight cubature for functions in Sobolev spaces. Furthermore, so-called sequences of Quasi-Monte Carlo (QMC) designs for certain Sobolev spaces have

[☆] Communicated by A. Hinrichs.

* Corresponding author.

E-mail addresses: anna.breger@univie.ac.at (A. Breger), martin.ehler@univie.ac.at (M. Ehler), mgraef@kfs.oeaw.ac.at (M. Gräf).

the optimal covering property, i.e., the underlying sequence of point sets has asymptotically optimal covering radii.

Here, we extend these results from sequences of QMC designs for Sobolev spaces on the sphere to sequences of weighted QMC designs for Sobolev spaces on compact smooth Riemannian manifolds. These theoretical results are derived by combining the ideas in [8] with the general findings on cubature points on compact smooth Riemannian manifolds in [7]. In the second part of the present manuscript, we numerically construct a sequence of QMC designs for Sobolev spaces on the Grassmannian manifold and illustrate numerically that their covering radii indeed behave in accordance with the theoretical results.

The special classes of cubature points in Grassmannians as studied in [2–5] from a theoretical point of view yield examples of sequences of (weighted) QMC designs for Sobolev spaces, see [19] for the construction through numerical minimization. For related results on cubature points in more classical settings, see [18,25,26,28,32,35,37] and, for further results, see [9,14,15,21,22,24,29–31] and references therein.

The outline of our paper is as follows. In Section 2, we briefly discuss the concept of asymptotically optimal covering radii. Low-cardinality weighted design sequences are introduced in Section 3, where we also state our main theoretical result. In Section 4, we state the bound on the covering radius by the worst case error of integration for Sobolev spaces. Section 5 is dedicated to the proof of this bound. In Section 7, we illustrate our theoretical results by numerical experiments for the Grassmannian manifold $\mathcal{G}_{2,4}$.

2. Optimal asymptotics of the covering radius

Let $\mathcal{M}$ be a compact smooth d -dimensional Riemannian manifold. We denote its normalized Riemannian measure by μ and its Riemannian distance by dist . Given any finite collection of points $\{x_j\}_{j=1}^n \subset \mathcal{M}$, the *covering radius* ρ is

$$\rho := \rho(\{x_j\}_{j=1}^n) := \sup_{x \in \mathcal{M}} \min_{1 \leq j \leq n} \text{dist}(x, x_j).$$

We denote the ball of radius r centered at $x \in \mathcal{M}$ by $B_r(x) := \{y \in \mathcal{M} : \text{dist}(x, y) \leq r\}$. The covering radius ρ of $\{x_j\}_{j=1}^n$ satisfies $\mathcal{M} = \bigcup_{j=1}^n B_\rho(x_j)$. By compactness of $\mathcal{M}$ we have¹

$$\mu(B_r(x)) \asymp r^d, \quad x \in \mathcal{M}, \quad 0 < r \leq 1, \quad (1)$$

where the constants neither depend on x nor r . Hence, the chain of inequalities $\mu(\mathcal{M}) \leq \sum_{j=1}^n \mu(B_\rho(x_j)) \lesssim n\rho^d$ yields the lower bound

$$n^{-\frac{1}{d}} \lesssim \rho. \quad (2)$$

Note that we shall keep using the notation $\{x_j\}_{j=1}^n$ throughout the manuscript but allow for multiple copies of the same vector.

Definition 2.1. Given a sequence of n_i -point sets $\{x_j^i\}_{j=1}^{n_i} \subset \mathcal{M}$, $i = 1, 2, \dots$, with $n_i \rightarrow \infty$, we say that the associated sequence of covering radii $(\rho_i)_{i=1}^\infty$ is *asymptotically optimal* if the lower bound (2) is matched, i.e., if $\rho_i \asymp n_i^{-\frac{1}{d}}$ and the constants do not depend on i .

According to [33], the expectation of the covering radius ρ of n random points, independently distributed according to μ , satisfies

$$\mathbb{E}\rho \asymp n^{-\frac{1}{d}} \log(n)^{\frac{1}{d}}, \quad (3)$$

¹ We use the symbols $\lesssim$ and $\gtrsim$ to indicate that the corresponding inequalities hold up to a positive constant factor on the respective right-hand side. The notation $\asymp$ means that both relations $\lesssim$ and $\gtrsim$ hold. If not explicitly stated, the dependence of the constants on other parameters shall be clear from the context.

where the constant is independent of n but may depend on $\mathcal{M}$ (and hence on d). Thus, independent uniformly distributed random points have optimal expected covering radius up to a logarithmic factor $\log(n)^{1/d}$. It should be noted that the asymptotic behavior in (3) for the sphere was already conjectured in [10] and then proved in [33].

3. Optimal coverings from low-cardinality cubatures

Let $\{\varphi_\ell\}_{\ell=0}^\infty$ be the collection of orthonormal eigenfunctions of the Laplace–Beltrami operator Δ on $\mathcal{M}$ with eigenvalues $\{-\lambda_\ell\}_{\ell=0}^\infty$ arranged by $0 = \lambda_0 \leq \lambda_1 \leq \dots$. Without loss of generality, we choose each φ_ℓ to be real-valued, in particular, $\varphi_0 \equiv 1$. We denote by $L_p(\mathcal{M})$, $1 \leq p < \infty$, the Banach space of complex-valued μ -measurable functions on $\mathcal{M}$, whose p th power of the absolute value is integrable (with the standard modifications when $p = \infty$).

The space of *diffusion polynomials* of bandwidth $t \geq 0$ is

$$\Pi_t := \text{span}\{\varphi_\ell : \lambda_\ell \leq t^2\}. \quad (4)$$

Given a set of points $\{x_j\}_{j=1}^n \subset \mathcal{M}$ and a set of associated weights $\{\omega_j\}_{j=1}^n \subset \mathbb{R}$, we say that $\{(x_j, \omega_j)\}_{j=1}^n$ is a *weighted t -design* if

$$\int_{\mathcal{M}} f(x) d\mu(x) = \sum_{j=1}^n \omega_j f(x_j) \quad \text{for all } f \in \Pi_t. \quad (5)$$

If (5) is satisfied with weights $\omega_j = 1/n$ for $j = 1, \dots, n$, then we call $\{x_j\}_{j=1}^n$ a *t -design*.

Weyl's estimates on the spectrum of an elliptic operator yield

$$\dim(\Pi_t) \asymp t^d,$$

cf. [27, Theorem 17.5.3]. Therefore, any sequence of weighted t_i -designs $\{(x_j^i, \omega_j^i)\}_{j=1}^{n_i}$ necessarily satisfies $n_i \gtrsim t_i^d$ as $t_i \rightarrow \infty$.

Definition 3.1. A sequence of (weighted) t_i -designs $\{(x_j^i, \omega_j^i)\}_{j=1}^{n_i}$, $i = 1, 2, \dots$ with $t_i \rightarrow \infty$ is called a *low-cardinality (weighted) design sequence* if

$$n_i \asymp t_i^d \quad (6)$$

and the constants are independent of i .

Note that there exist sequences of weighted t_i -designs $\{(x_j^i, \omega_j^i)\}_{j=1}^{n_i}$, $i = 1, 2, \dots$, with positive weights satisfying (6), cf. [17]. We now state our main theoretical result.

Theorem 3.2. *The sequence of covering radii of any low-cardinality weighted design sequence with positive weights is asymptotically optimal.*

The remaining part of the present paper is dedicated to prove Theorem 3.2 and to numerically illustrate our results on the Grassmannian manifold. We conclude this section with a remark concerning the weights.

Remark 3.3. There exist low-cardinality design sequences for the unit sphere $\mathcal{M} = \mathbb{S}^d \subset \mathbb{R}^{d+1}$ and the Grassmannian manifold, cf. [6,20]. For general $\mathcal{M}$, however, we only know that t -designs with arbitrarily large t exist, cf. [34], but it is still an open problem whether or not the asymptotics (6) can be achieved.

4. Worst case error of integration

To prove Theorem 3.2, we follow the approach for the sphere in [8]. We shall first introduce the worst case error of a cubature method for Sobolev spaces and shall check that it provides an

upper bound on the covering radius. Next, we consider the concept of sequences of weighted QMC designs for Sobolev spaces, analogous to [11], and verify that the sequence of associated covering radii is asymptotically optimal. It follows from [7] that low-cardinality weighted design sequences with positive weights are indeed sequences of weighted QMC designs for Sobolev spaces.

The *worst case error* of a cubature method with integration nodes $x_1, \dots, x_n \in \mathcal{M}$ and associated weights $\omega_1, \dots, \omega_n \in \mathbb{R}$ for functions in the unit ball of a Banach space $\mathcal{H}$ of complex-valued functions on $\mathcal{M}$ is

$$\text{wce}(\{(x_j, \omega_j)\}_{j=1}^n, \mathcal{H}) := \sup_{\substack{f \in \mathcal{H} \\ \|f\|_{\mathcal{H}} \leq 1}} \left| \int_{\mathcal{M}} f(x) d\mu(x) - \sum_{j=1}^n \omega_j f(x_j) \right|. \quad (7)$$

Although suppressed by our notation, (7) depends on the particular norm $\|\cdot\|_{\mathcal{H}}$, which shall always be clear from the context in the present manuscript. For most parts, we take $\mathcal{H}$ to be a Sobolev space, which we define next. The Fourier transform of $f \in L_p(\mathcal{M})$ with $1 \leq p \leq \infty$ is

$$\hat{f}(\ell) := \int_{\mathcal{M}} f(x) \varphi_{\ell}(x) d\mu(x), \quad \ell = 0, 1, 2, \dots,$$

and extends to distributions on $\mathcal{M}$. The Sobolev space $H_p^s(\mathcal{M})$ for $1 \leq p \leq \infty$ and $s > 0$ is the set of all distributions on $\mathcal{M}$ with $(I - \Delta)^{s/2} f \in L_p(\mathcal{M})$, i.e., with

$$\|f\|_{H_p^s} := \|(I - \Delta)^{s/2} f\|_{L_p} = \left\| \sum_{\ell=0}^{\infty} (1 + \lambda_{\ell})^{s/2} \hat{f}(\ell) \varphi_{\ell} \right\|_{L_p} < \infty. \quad (8)$$

Note that $H_p^s(\mathcal{M})$ is contained in the space of continuous functions on $\mathcal{M}$ provided that $s > d/p$, cf. [7] and [36, Theorem 7.4.5, Section 7.4.2]. We shall stick to this range throughout the present paper.

It turns out that the covering radius is bounded by the worst case error. The following result has been derived in [8] for $\mathcal{M}$ being the sphere $\mathbb{S}^d$ and constant weights $\omega_j = 1/n$.

Proposition 4.1. *Let $1 \leq p \leq \infty$ and $s > d/p$ be given. Then for any n -point set $\{x_j\}_{j=1}^n \subset \mathcal{M}$ and any associated weight set $\{\omega_j\}_{j=1}^n \subset \mathbb{R}$ the covering radius ρ of $\{x_j\}_{j=1}^n$ satisfies*

$$\rho \lesssim [\text{wce}(\{(x_j, \omega_j)\}_{j=1}^n, H_p^s(\mathcal{M}))]^{1/(s+d/q)}, \quad (9)$$

where $1/p + 1/q = 1$. The constant may depend on $\mathcal{M}$, s , and p .

We shall postpone the proof of Proposition 4.1 to Section 5. In the following we turn to sequences of points whose worst case error of integration satisfies suitable decay conditions. This decay connects to the covering radius via the bound (9). The case of the sphere with constant weights in the following definition is due to [11]:

Definition 4.2. Given $1 \leq p \leq \infty$ and $s > d/p$, a sequence $\{(x_j^i, \omega_j^i)\}_{j=1}^{n_i}$, $i = 1, 2, \dots$, of n_i -point sets $\{x_j^i\}_{j=1}^{n_i} \subset \mathcal{M}$ and n_i -weight sets $\{\omega_j^i\}_{j=1}^{n_i} \subset \mathbb{R}$ with $n_i \rightarrow \infty$ is called a *sequence of weighted QMC designs* for $H_p^s(\mathcal{M})$ if

$$\text{wce}(\{(x_j^i, \omega_j^i)\}_{j=1}^{n_i}, H_p^s(\mathcal{M})) \lesssim n_i^{-\frac{s}{d}}. \quad (10)$$

According to [7], low-cardinality weighted design sequences with positive weights satisfy the conditions in Definition 4.2:

Proposition 4.3 ([7]). *Any low-cardinality weighted design sequence with positive weights is a sequence of weighted QMC designs for $H_p^s(\mathcal{M})$ with $s > d/p$ and $1 \leq p \leq \infty$.*

Due to (8), for $1 \leq p \leq p' \leq \infty$, the space $H_{p'}^s(\mathcal{M})$ is continuously embedded into $H_p^s(\mathcal{M})$ with $\|f\|_{H_p^s} \leq \|f\|_{H_{p'}^s}$ for $f \in H_{p'}^s(\mathcal{M})$. Thus, any sequence of weighted QMC designs for $H_{p'}^s(\mathcal{M})$ with $s > d/p$ is also a sequence of weighted QMC designs for $H_p^s(\mathcal{M})$. Therefore, $p = 1$ is the strongest requirement. Note that, for the sphere, the above embedding properties have also been stated in [8, Theorem 4.2].

Remark 4.4. If $\{(x_j^i, \omega_j^i)\}_{j=1}^{n_i}$ is a sequence of weighted QMC designs for $H_p^s(\mathcal{M})$ for some $p \geq 1$, then Proposition 4.1 yields that the covering radius ρ_i of $\{x_j^i\}_{j=1}^{n_i}$ is bounded by

$$\rho_i \lesssim n_i^{-\frac{1}{d}(\frac{s}{s+d(1-1/p)})}.$$

Thus, any sequence of weighted QMC designs for $H_1^s(\mathcal{M})$ provides

$$\rho_i \lesssim [\text{wce}(\{(x_j^i, \omega_j^i)\}_{j=1}^{n_i}, H_1^s(\mathcal{M}))]^{1/s} \lesssim n_i^{-\frac{1}{d}},$$

so that Theorem 3.2 is a consequence of Propositions 4.1 and 4.3. Hence, to complete the proof of Theorem 3.2, it only remains to verify Proposition 4.1, which is the topic of the subsequent section.

5. Proof of Proposition 4.1

The proof of Proposition 4.1 relies on findings in [7]. We recapitulate the following localization result, which is one of the key ingredients.

Lemma 5.1. For $s \in \mathbb{R}$ and $h : \mathbb{R} \rightarrow \mathbb{R}$ a smooth function supported on the interval $[\frac{1}{2}, 2]$, the kernel

$$K_r^s(x, y) := \sum_{\ell=0}^{\infty} h\left(\frac{\lambda_{\ell}}{r}\right) (1 + \lambda_{\ell})^{-s/2} \varphi_{\ell}(x) \varphi_{\ell}(y) \quad (11)$$

is bounded by

$$\left\| \int_{\mathcal{M}} |K_r^s(\cdot, y)| d\mu(y) \right\|_{L_{\infty}} \lesssim r^{-s}, \quad r > 0, \quad (12)$$

where the constant does not depend on r .

Proof. According to [7, Lemma 2.8], the estimate

$$|K_r^s(x, y)| \lesssim r^{d-s} (1 + r \text{dist}(x, y))^{-(d+1)}, \quad r > 0,$$

holds and leads to the desired assertion

$$\begin{aligned} \sup_{x \in \mathcal{M}} \int_{\mathcal{M}} |K_r^s(x, y)| d\mu(y) &\lesssim r^{d-s} \sup_{x \in \mathcal{M}} \int_{\mathcal{M}} (1 + r \text{dist}(x, y))^{-(d+1)} d\mu(y) \\ &\lesssim r^{-s} \sup_{x \in \mathcal{M}} \int_{\mathcal{M}} r^d (1 + r \text{dist}(x, y))^{-(d+1)} d\mu(y) \\ &\lesssim r^{-s}. \quad \square \end{aligned}$$

We shall make use of Lemma 5.1 to verify the following result.

Lemma 5.2. Let $R, s > 0$, and $1 \leq p \leq \infty$ be fixed. For any $0 < r \leq R$ and $z \in \mathcal{M}$, there is a function $f := f_{r,z} \in C^{\infty}(\mathcal{M})$ with support in $B_r(z)$ such that

$$\|f\|_{H_p^s} \lesssim r^{-s+d/p}, \quad r^d \lesssim \int_{\mathcal{M}} f(x) d\mu(x), \quad (13)$$

where the constants neither depend on z nor r .

Proof. As in [7, Proof of Theorem 2.16], for any radius $0 < r \leq R$ and $z \in \mathcal{M}$, there is a function $f := f_{r,z} \in C^{\infty}(\mathcal{M})$ with support in $B_r(z)$ such that

$$\|\Delta^{\ell} f\|_{L_{\infty}} \lesssim r^{-2\ell} \quad \text{for } \ell = 0, 1, 2, \dots, \quad r^d \lesssim \int_{\mathcal{M}} f(x) d\mu(x), \quad (14)$$

where by compactness of $\mathcal{M}$ the constants neither depend on z nor r .

Similarly as in [7, Proof of Theorem 2.16], we bound the norm $\|f\|_{H_p^s} = \|(I - \Delta)^{s/2}f\|_{L_p}$ by using a dyadic decomposition of unity of the Fourier domain, see, for instance, [37]. That is, we set $h(x) := g(x) - g(2x)$ for some smooth function $g : \mathbb{R} \rightarrow \mathbb{R}$ satisfying

$$g(x) = \begin{cases} 1, & x \leq 1, \\ 0, & x \geq 2, \end{cases}$$

and obtain $\text{supp}(h) \subset [1/2, 2]$ with

$$\sum_{m=-\infty}^{\infty} h(2^{-m}x) = \begin{cases} 1, & x > 0, \\ 0, & \text{else.} \end{cases}$$

Using the Fourier expansion $f = \sum_{\ell=0}^{\infty} \hat{f}(\ell)\varphi_{\ell}$ we arrive at

$$(I - \Delta)^{s/2}f = \hat{f}(0) + \sum_{\ell=0}^{\infty} \sum_{m=-\infty}^{\infty} h(2^{-m}\lambda_{\ell})(1 + \lambda_{\ell})^{s/2}\hat{f}(\ell)\varphi_{\ell}.$$

Fixing an integer $L > s/2$ and applying (11) with $\hat{f}(\ell) = \int_{\mathcal{M}} f(y)\varphi_{\ell}(y)d\mu(y)$ lead to

$$\begin{aligned} (I - \Delta)^{s/2}f &= \hat{f}(0) + \sum_{2^m \leq 1/r} \int_{\mathcal{M}} K_{2^m}^{-s}(\cdot, y)f(y)d\mu(y) \\ &\quad + \sum_{2^m > 1/r} \int_{\mathcal{M}} K_{2^m}^{2L-s}(\cdot, y)(I - \Delta)^L f(y)d\mu(y). \end{aligned}$$

We shall now bound the three terms of the right hand side separately. First, the Hölder inequality with (14) for $\ell = 0$ yields

$$|\hat{f}(0)| = \left| \int_{\mathcal{M}} f(y)d\mu(y) \right| \leq \|f\|_{L_{\infty}}\mu(\mathcal{M}) \lesssim 1 \lesssim r^{-s},$$

where the last estimate is due to r being bounded from above and $s > 0$. To bound the second term, we apply Lemma 5.1 and derive

$$\begin{aligned} \left\| \sum_{2^m \leq 1/r} \int_{\mathcal{M}} K_{2^m}^{-s}(\cdot, y)f(y)d\mu(y) \right\|_{L_{\infty}} &\leq \sum_{2^m \leq 1/r} \left\| \int_{\mathcal{M}} |K_{2^m}^{-s}(\cdot, y)|d\mu(y) \right\|_{L_{\infty}} \|f\|_{L_{\infty}} \\ &\lesssim \sum_{2^m \leq 1/r} \left\| \int_{\mathcal{M}} |K_{2^m}^{-s}(\cdot, y)|d\mu(y) \right\|_{L_{\infty}} \\ &\lesssim \sum_{2^m \leq 1/r} 2^{ms} \lesssim r^{-s}. \end{aligned}$$

Lemma 5.1 also leads to a bound on the third term by means of

$$\begin{aligned} \left\| \sum_{2^m > 1/r} \int_{\mathcal{M}} K_{2^m}^{2L-s}(\cdot, y)(I - \Delta)^L f(y)d\mu(y) \right\|_{L_{\infty}} \\ &\lesssim \sum_{2^m > 1/r} \left\| \int_{\mathcal{M}} K_{2^m}^{2L-s}(\cdot, y)(I + \Delta)^L f(y)d\mu(y) \right\|_{L_{\infty}} \\ &\lesssim \sum_{2^m > 1/r} \left\| \int_{\mathcal{M}} K_{2^m}^{2L-s}(\cdot, y)d\mu(y) \right\|_{L_{\infty}} \|(I - \Delta)^L f\|_{L_{\infty}} \end{aligned}$$

$$\begin{aligned} &\lesssim \sum_{2^m > 1/r} 2^{-(2L-s)m} r^{-2L} \\ &\lesssim r^{2L-s} r^{-2L} = r^{-s}, \end{aligned}$$

where we have also applied (14) and r being bounded from above.

Thus, we obtain

$$\|(I - \Delta)^{s/2} f\|_{L_\infty} \lesssim r^{-s}, \quad (15)$$

where the constants neither depend on z nor r . To cover the range $1 \leq p < \infty$, we recall that f is supported on $B_r(z)$, so that the Hölder inequality with (1) yields

$$\|(I - \Delta)^{s/2} f\|_{L_p}^p \lesssim r^{-ps} \mu(B_r(z)) \lesssim r^{-sp+d}, \quad (16)$$

which concludes the proof. $\square$

We are now prepared to complete the proof of our main theoretical result.

Proof of Proposition 4.1. For a given point set $\{x_j\}_{j=1}^n \subset \mathcal{M}$ with covering radius ρ , let z be a center of a maximal hole, i.e.,

$$\mathring{B}_\rho(z) \cap \{x_j\}_{j=1}^n = \emptyset, \quad (17)$$

where $\mathring{B}_\rho(z) = \{x \in \mathcal{M} : \text{dist}(x, z) < \rho\}$ denotes the interior of $B_\rho(z)$. Note that ρ is bounded by the diameter of $\mathcal{M}$. Let $f = f_{\rho,z} \in \mathcal{C}^\infty(\mathcal{M})$ be as in Lemma 5.2, i.e., $\text{supp}(f) \subset B_\rho(z)$ and (13) holds with $r = \rho$. Since f must vanish outside of $\mathring{B}_\rho(z)$, (17) implies $f(x_j) = 0$ for all $j = 1, \dots, n$. Thus, the definition of the worst case error of integration yields

$$\text{wce}(\{(x_j, \omega_j)\}_{j=1}^n, H_p^s(\mathcal{M})) \geq \frac{|\int_{\mathcal{M}} f(x) d\mu(x) - 0|}{\|f\|_{H_p^s}} \gtrsim \frac{\rho^d}{\rho^{-s+d/p}} = \rho^{s+d/q},$$

with $1/p + 1/q = 1$, and the constant neither depends on z nor ρ . $\square$

6. Worst case error by means of reproducing kernels

Before we can provide numerical experiments, we need to discuss a strategy to compute the worst case error of a cubature method in the Hilbert space $H_2^s(\mathcal{M})$ with inner product

$$\langle f, g \rangle_{H_2^s} = \sum_{\ell=0}^{\infty} (1 + \lambda_\ell)^s \hat{f}(\ell) \overline{\hat{g}(\ell)}, \quad f, g \in H_2^s(\mathcal{M}). \quad (18)$$

For $s > d/2$, the Bessel kernel $K_B^s : \mathcal{M} \times \mathcal{M} \rightarrow \mathbb{R}$ defined by

$$K_B^s(x, y) = \sum_{\ell=0}^{\infty} (1 + \lambda_\ell)^{-s} \varphi_\ell(x) \varphi_\ell(y) \quad (19)$$

is the reproducing kernel for $H_2^s(\mathcal{M})$ and induces the inner product (18) and the associated norm. For later reference, we shall express the worst case error by means of the reproducing kernel. If $K : \mathcal{M} \times \mathcal{M} \rightarrow \mathbb{R}$ is a reproducing kernel for some reproducing kernel Hilbert space H_K of continuous functions on $\mathcal{M}$, then the worst case error of integration is

$$\begin{aligned} \text{wce}(\{(x_j, \omega_j)\}_{j=1}^n, H_K)^2 &= \sum_{j,j'=1}^n \omega_j \omega_{j'} K(x_j, x_{j'}) - 2 \sum_{j=1}^n \omega_j \int_{\mathcal{M}} K(x_j, x) d\mu(x) \\ &\quad + \int_{\mathcal{M}} \int_{\mathcal{M}} K(x, y) d\mu(x) d\mu(y), \end{aligned} \quad (20)$$

cf. [23, Theorem 2.7], see also [32]. Note that the norm in H_K is uniquely determined by its reproducing kernel K . If K has the Fourier expansion

$$K(x, y) = \sum_{\ell=0}^{\infty} \hat{K}(\ell) \varphi_{\ell}(x) \varphi_{\ell}(y) \quad (21)$$

with $\hat{K}(\ell) \in \mathbb{C}$ for $\ell = 0, 1, \dots$, then (20) becomes

$$\text{wce}(\{(x_j, \omega_j)\}_{j=1}^n, H_K)^2 = \sum_{j,j'=1}^n \omega_j \omega_{j'} K(x_j, x_{j'}) + \hat{K}(0) \left(1 - 2 \sum_{j=1}^n \omega_j\right), \quad (22)$$

where we have used $\varphi_0 \equiv 1$. For the Bessel kernel K_B^s , we observe $\hat{K}_B^s(\ell) = (1 + \lambda_{\ell})^{-s}$ with $\hat{K}_B^s(0) = 1$

We conclude this section by the following result on the worst case error of uniformly distributed random points, for which constant weights $\omega_j = 1/n$ are the natural choice.

Proposition 6.1. *If K is a reproducing kernel on $\mathcal{M}$ and $x_1, \dots, x_n$ are random points on $\mathcal{M}$, independently distributed according to μ , then it holds*

$$\sqrt{\mathbb{E} \left[\text{wce}(\{(x_j, \frac{1}{n})\}_{j=1}^n, H_K)^2 \right]} = cn^{-\frac{1}{2}}, \quad (23)$$

where

$$c^2 = \int_{\mathcal{M}} K(x, x) d\mu(x) - \int_{\mathcal{M}} \int_{\mathcal{M}} K(x, y) d\mu(x) d\mu(y). \quad (24)$$

By applying (21), the constant in (24) is $c^2 = \sum_{\ell=1}^{\infty} \hat{K}(\ell)$. For the Bessel kernel K_B^s , the condition $s > d/2$ implies that $\frac{s}{d} > \frac{1}{2}$, so that on average sequences of weighted QMC designs for Sobolev spaces indeed perform better than random points for sufficiently smooth functions. For the proof of Proposition 6.1, we refer to [23, proof of Corollary 2.8], see also [32]. Eq. (23) is derived in [11] for the sphere.

7. Numerical experiments for the Grassmannian manifold

This section provides numerical results that illustrate Proposition 4.3 and Theorem 3.2 for the case of the Grassmannian manifold, i.e., the collection of k -dimensional linear subspaces in $\mathbb{R}^m$. We identify the Grassmannian with the set of orthogonal projectors on $\mathbb{R}^m$ of rank k , denoted by

$$\mathcal{G}_{k,m} := \{P \in \mathbb{R}^{m \times m} : P^{\top} = P, P^2 = P, \text{trace}(P) = k\}. \quad (25)$$

Hence, the Grassmannian $\mathcal{G}_{k,m}$ is a $d = k(m - k)$ -dimensional submanifold of the Euclidean space $\mathbb{R}^{m^2} \cong \mathbb{R}^{m \times m}$. Moreover $\mathbb{R}^{m^2}$ induces a Riemannian metric, which in turn yields the canonical probability measure and the canonical geodesic distance on the Grassmannian $\mathcal{G}_{k,m}$ denoted by $\mu_{k,m}$ and $\text{dist}_{k,m}$, respectively. In particular, the geodesic distance $\text{dist}_{k,m}(P, Q)$ between $P, Q \in \mathcal{G}_{k,m}$ is proportional to the 2-norm of the corresponding principal angles $\theta_1, \dots, \theta_k$ between the subspaces associated to P and Q , respectively. More precisely, the distance can be computed by

$$\text{dist}_{k,m}(P, Q) = \sqrt{2} \sqrt{\theta_1^2 + \dots + \theta_k^2}, \quad (26)$$

where $\theta_i = \arccos(\sqrt{y_i})$ and $y_1, \dots, y_k$ are the k -largest eigenvalues of PQ (counted with multiplicities). Note that the factor $\sqrt{2}$ is often omitted in the literature but accounts for the particular embedding (25) since then

$$\text{dist}_{k,m}(P, Q) = \|P - Q\|_F + o(\|P - Q\|_F^2), \quad P, Q \in \mathcal{G}_{k,m}, \quad (27)$$

where $\|X\|_F$ is the Frobenius-norm of $X \in \mathbb{R}^{m \times m}$.

7.1. Numerical construction of t -designs

Theorem 3.2 tells us that the sequence of covering radii of a low-cardinality design sequence is asymptotically optimal. To illustrate this result, we first construct a sequence of t_i -designs. It is known that any collection of points $\{P_j\}_{j=1}^n \subset \mathcal{G}_{k,m}$ satisfies

$$\frac{1}{n^2} \sum_{j,j'=1}^n \text{trace}(P_j, P_{j'})^i \geq \int_{\mathcal{G}_{k,m}} \int_{\mathcal{G}_{k,m}} \text{trace}(P, Q)^i d\mu_{k,m}(P) d\mu_{k,m}(Q) \quad (28)$$

for $i = 1, 2, \dots$, cf. [5]. According to [12], equality in (28) yields a t_i -design with $t_i = i \frac{2}{\sqrt{k}}$, see also [19]. Hence, (28) enables a simple approach to numerically compute designs by minimization and checking for equality.

Our numerical experiments shall focus on $\mathcal{G}_{2,4}$, which has dimension $d = 2(4 - 2) = 4$, so that for low-cardinality design sequences the cardinality n_i of its t_i -designs must satisfy $n_i \asymp i^4 \asymp t_i^4$, see (6). Indeed, we have chosen

$$n_i = \left\lfloor \frac{1}{3}(i+1)^2(1+i+\frac{1}{2}i^2) \right\rfloor$$

and computed points $\{P_j^i\}_{j=1}^{n_i} \subset \mathcal{G}_{2,4}$, for $i = 1, \dots, 14$, by a nonlinear conjugate gradient method on manifolds, cf. [1,12,23], such that

$$\left| \int_{\mathcal{G}_{2,4}} f(P) d\mu_{2,4}(P) - \frac{1}{n_i} \sum_{j=1}^{n_i} f(P_j^i) \right| < 10^{-7} \quad (29)$$

for all $f \in \Pi_{t_i}$ with $\|f\|_{L_2} \leq 1$. Although the left-hand side of (29) may not be zero exactly, we shall refer to $\{P_j^i\}_{j=1}^{n_i}$ as t_i -designs in the following.

7.2. Integration

In view of bounding the covering radius, we would like to provide numerical experiments on the worst case error of cubature methods for Sobolev spaces $H_p^s(\mathcal{G}_{k,m})$ for $p = 1$. However, determining the worst case error is a tough task in general. For $p = 2$ and $s > k(m - k)/2$, on the other hand, we are dealing with reproducing kernel Hilbert spaces, in which case we can invoke (22) provided that its reproducing kernel is numerically accessible.

Our first numerical experiments are about integration in $H_2^{7/2}(\mathcal{G}_{2,4})$, so that the worst case error is indeed given by (22). However, the infinite series of the Bessel kernel K_B^s in (19) is numerically cumbersome, so that we consider the positive definite kernel

$$K_1(P, Q) = k_1(\text{trace}(PQ)), \quad P, Q \in \mathcal{G}_{2,4},$$

where

$$k_1(r) = \sqrt{(2-r)^3} + 2r, \quad 0 \leq r \leq 2.$$

According to [12], see also [7], its reproducing kernel Hilbert space H_{K_1} is the Sobolev space $H_2^{7/2}(\mathcal{G}_{2,4})$ equipped with an equivalent norm, i.e., the two norms $\|\cdot\|_{H_2^{7/2}}$ and $\|\cdot\|_{K_1}$ are comparable. This implies

$$\text{wce}(\{(P_j, \frac{1}{n})\}_{j=1}^n, H_2^{7/2}(\mathcal{G}_{2,4})) \asymp \text{wce}(\{(P_j, \frac{1}{n})\}_{j=1}^n, H_{K_1}),$$

where the constants are independent of the particular point set $\{P_j\}_{j=1}^n \subset \mathcal{G}_{2,4}$. According to Proposition 4.3, the error $\text{wce}(\{(P_j^i, \frac{1}{n_i})\}_{j=1}^{n_i}, H_{K_1})$ decays like $n_i^{-7/8}$ for low-cardinality design sequences. Hence, we expect a linear behavior with slope $-7/8$ in logarithmic error plots.

In a second numerical experiment on integration, we shall consider a second positive definite kernel K_2 given by

$$K_2(P, Q) = k_2(\text{trace}(PQ)), \quad P, Q \in \mathcal{G}_{2,4},$$

where

$$k_2(r) = \frac{3}{2} \exp(-(2-r)), \quad 0 \leq r \leq 2.$$

Its reproducing kernel Hilbert space H_{K_2} satisfies

$$H_{K_2} \subset \bigcap_{s>2} H_2^s(\mathcal{G}_{2,4}). \quad (30)$$

According to [Proposition 4.3](#), a low-cardinality design sequence is a sequence of QMC designs for $H_2^s(\mathcal{G}_{2,4})$ for all $s > 2$. Due to (30), we expect a super linear behavior of $\text{wce}(\{(P_j^i, \frac{1}{n_i})\}_{j=1}^{n_i}, H_{K_2})$ in logarithmic plots.

We now further specify the worst case error in H_{K_1} and H_{K_2} via (22).

Lemma 7.1. *The worst case errors in H_{K_1} and H_{K_2} are*

$$\text{wce}(\{(P_j, \frac{1}{n})\}_{j=1}^n, H_{K_1})^2 = \frac{1}{n^2} \sum_{j,j'=1}^n K_1(P_j, P_{j'}) \quad (31)$$

$$- (2 + \frac{74}{75}\sqrt{2} - \frac{2}{5} \log(1 + \sqrt{2})),$$

$$\text{wce}(\{(P_j, \frac{1}{n})\}_{j=1}^n, H_{K_2})^2 = \frac{1}{n^2} \sum_{j,j'=1}^n K_2(P_j, P_{j'}) - \frac{3}{2} \exp(-1) \text{Shi}(1), \quad (32)$$

respectively, where $\text{Shi}(x) = \int_0^x \frac{\sinh(t)}{t} dt$ is the hyperbolic sine integral.

Proof. In view of (22) it remains to compute the 0-th Fourier coefficients

$$\hat{K}_i(0) = \int_{G_{2,4}} K_i(P, Q) d\mu_{2,4}(P) d\mu_{2,4}(Q), \quad i = 1, 2.$$

According to [16, Example 4.3], the orthogonal invariance of K_1 and K_2 with the variable transformation $\xi_{\pm} = \cos(\theta_1 \pm \theta_2)$, where θ_1, θ_2 are the principal angles between P and Q , yield

$$\int_{G_{2,4}} K_i(P, Q) d\mu_{2,4}(P) d\mu_{2,4}(Q) = \int_{-1}^1 \int_{|\xi_-|}^1 k_i(1 + \xi_- \xi_+) d\xi_+ d\xi_-, \quad i = 1, 2.$$

The symmetry of the function $(\xi_-, \xi_+) \mapsto \xi_- \xi_+$ leads to

$$\hat{K}_i(0) = \frac{1}{4} \int_{-1}^1 \int_{-1}^1 k_i(1 + \xi_- \xi_+) d\xi_+ d\xi_-, \quad i = 1, 2.$$

For $i = 1$, this yields

$$\begin{aligned} \hat{K}_1(0) &= \frac{1}{4} \int_{-1}^1 \int_{-1}^1 ((1 - \xi_- \xi_+)^{\frac{3}{2}} + 2\xi_- \xi_+ + 2) d\xi_+ d\xi_- \\ &= 2 + \int_{-1}^1 \frac{(1+\xi_-)^{\frac{5}{2}} - (1-\xi_-)^{\frac{5}{2}}}{10\xi_-} d\xi_-. \end{aligned}$$

The assertion (31) is then checked by a computer algebra system.

For $i = 2$, we obtain

$$\begin{aligned} \hat{K}_2(0) &= \frac{3}{8} \int_{-1}^1 \int_{-1}^1 \exp(\xi_- \xi_+ - 1) d\xi_+ d\xi_- \\ &= \frac{3}{8} \exp(-1) \int_{-1}^1 \frac{\exp(\xi_-) - \exp(-\xi_-)}{\xi_-} d\xi_-, \end{aligned}$$

so that (32) follows immediately. $\square$

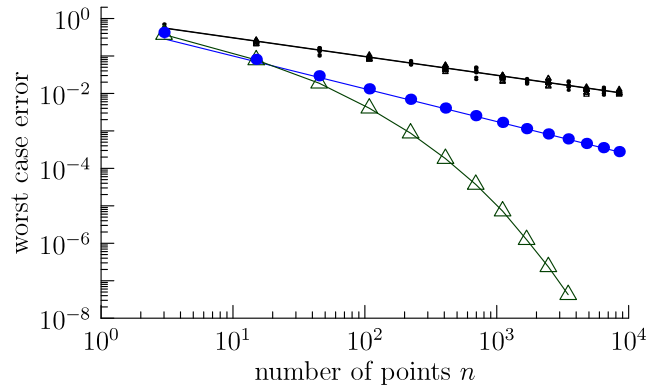


Fig. 1. Integration by random sampling vs. t -designs. Dots refer to the worst case error in H_{K_1} , triangles to H_{K_2} . Big symbols correspond to t_i -designs, small symbols to n_i random points. The black line is the expectation of the worst case error and has slope $-1/2$, cf. Proposition 6.1, and the blue line's slope is $-7/8$. (For interpretation of the references to color in this figure legend, the reader is referred to the web version of this article.)

Fig. 1 shows logarithmic plots of

$$\text{wce}(\{(P_j^i, \frac{1}{n_i})\}_{j=1}^{n_i}, H_{K_1}), \quad \text{wce}(\{(P_j, \frac{1}{n_i})\}_{j=1}^{n_i}, H_{K_2})$$

for the t_i -designs $\{P_j^i\}_{j=1}^{n_i} \subset \mathcal{G}_{2,4}$, $i = 1, \dots, 14$, from Section 7.1. For comparison, it displays the worst case error of integration for random points, independently distributed according to $\mu_{2,4}$, cf. Proposition 6.1. Indeed, we observe the superior integration quality of the t_i -designs over random points. The theoretical results in Propositions 4.3 and 6.1 are in perfect accordance with the numerical experiment. The integration errors of the random points scatter around the expected integration error with rate $n_i^{-1/2}$ in both cases, H_{K_1} and H_{K_2} . The worst case error of low-cardinality design sequences achieves the optimal rate of $n_i^{-7/8}$ for functions in the Sobolev space $H_2^{7/2}(\mathcal{G}_{2,4})$. Due to $H_{K_2} \subset \bigcap_{s>2} H_2^s(\mathcal{G}_{2,4})$, we observe the super linear behavior in the logarithmic plots for $\text{wce}(\{(P_j^i, \frac{1}{n_i})\}_{j=1}^{n_i}, H_{K_2})$.

To conclude this section, we point out that Lemma 7.1 provides analytic expressions for the worst case errors in particular reproducing kernel Hilbert spaces. If we go beyond Hilbert spaces, then numerically computing worst case errors in Sobolev spaces becomes very challenging, which is illustrated by the following remark.

Remark 7.2. Let I_2 denote the 4 by 4 matrix with two ones in the left upper diagonal entries and zeros elsewhere. We shall consider the integration error with respect to a sequence of low-cardinality weighted design sequence $\{(P_j^i, \omega_j^i)\}_{j=1}^{n_i}$ for Π_{t_i} in $\mathcal{G}_{2,4}$. Without loss of generality we assume $P_1^i = I_2$, for $i = 1, 2, \dots$. The function $f_1(P) := K_1(I_2, P)$, for $P \in \mathcal{G}_{2,4}$, is contained in $H_2^{7/2}(\mathcal{G}_{2,4})$, so that its integration error decays at least as fast as $n^{-7/8}$. Numerical experiments suggest that it decays exactly at this rate, so that f_1 seems to be a single representative for the worst case error of integration in the reproducing kernel Hilbert space $H_2^{7/2}(\mathcal{G}_{2,4})$.

Now, we point out, although $f_1 \in H_\infty^3(\mathcal{G}_{2,4})$ with $f_1 \notin H_\infty^{3+\epsilon}(\mathcal{G}_{2,4})$, for all $\epsilon > 0$, its integration error decays faster than the worst case error in $H_\infty^3(\mathcal{G}_{2,4})$, which is just $n^{-3/4}$.

7.3. Approximating covering radii by random points

Given a sequence of n_i -point sets $\{P_j^i\}_{j=1}^{n_i} \subset \mathcal{G}_{k,m}$, for $i = 1, 2, \dots$, it is a tough task to numerically determine the exact covering radii ρ_i . In order to obtain a reasonable approximation, we generate $n \gg n_i$ independent random points $R_1, \dots, R_n \in \mathcal{G}_{k,m}$ and determine

$$\hat{\rho}_{i,n} := \max_{1 \leq j' \leq n} \min_{1 \leq j \leq n_i} \text{dist}_{k,m}(R_{j'}, P_j^i). \quad (33)$$

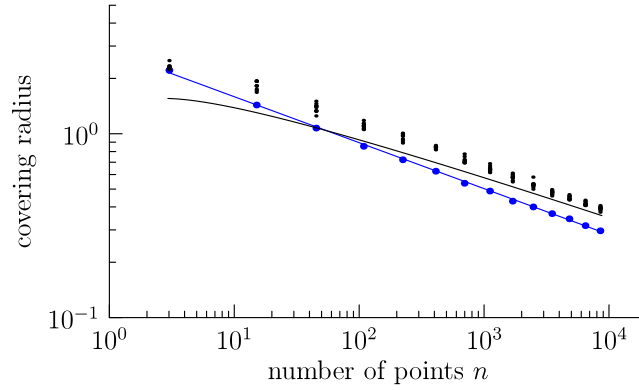


Fig. 2. Covering radii of random points vs. t -designs. Big dots correspond to t_i -designs, small dots to n_i random points. The blue line's slope is $-1/4$. The black curve corresponds to $2n^{-1/4} \log(n)^{1/4}$, which relates to the covering radii ϱ_n via (38).

Note that ρ_i is sandwiched by

$$\hat{\rho}_{i,n} \leq \rho_i \leq \hat{\rho}_{i,n} + \varrho_n, \quad (34)$$

where ϱ_n is the covering radius of $\{R_{j'}\}_{j'=1}^n$. Supposed that ϱ_n is sufficiently small, $\hat{\rho}_{i,n}$ is a decent approximation of ρ_i .

We aim to provide evidence that the computed approximation $\hat{\rho}_{i,n}$ is sufficiently accurate to numerically illustrate Theorem 3.2. To ballpark ϱ_n , recall that the covering radius ϱ_n of n random points, independently distributed according to $\mu_{k,m}$ on $\mathcal{G}_{k,m}$, behaves asymptotically as stated in (3) with $d = k(m - k)$. More precisely, [33, Corollary 3.3] yields together with the relation (27) that

$$\lim_{n \rightarrow \infty} \mathbb{E} \varrho_n \left(\frac{n}{\log(n)} \right)^{\frac{1}{k(m-k)}} = \left(\frac{\text{vol}(\mathcal{G}_{k,m})}{\text{vol}(\mathbb{B}_{k(m-k)})} \right)^{\frac{1}{k(m-k)}}, \quad (35)$$

where $\mathbb{B}_{k(m-k)}$ is the unit ball in $\mathbb{R}^{k(m-k)}$ and vol are the canonical volumes induced by the Euclidean metric. Thus, for large n , we expect that ϱ_n behaves like

$$\varrho_n \approx \left(\frac{\text{vol}(\mathcal{G}_{k,m})}{\text{vol}(\mathbb{B}_{k(m-k)})} \frac{n}{\log(n)} \right)^{\frac{1}{k(m-k)}}. \quad (36)$$

The volume of the Grassmannian is

$$\text{vol}(\mathcal{G}_{k,m}) = \frac{\Gamma_k(k/2)}{\Gamma_k(m/2)} (2\pi)^{k(m-k)/2}, \quad (37)$$

cf. [13, Section 1.4], where Γ_k is the multivariate Gamma function

$$\Gamma_k(x) = \pi^{k(k-1)/4} \prod_{i=1}^k \Gamma(x + (1-i)/2), \quad x > \frac{1}{2}(k-1).$$

Note, the formula (37) differs by a factor of $\sqrt{2}^{k(m-k)}$ from [13, Eq. (1.4.11)] due to our scaling with $\sqrt{2}$ in the geodesic distance, cf. (26). The volume of $\mathbb{B}_{k(m-k)}$ in $\mathbb{R}^{k(m-k)}$ is

$$\text{vol}(\mathbb{B}_{k(m-k)}) = \frac{\pi^{k(m-k)/2}}{\Gamma(k(m-k)/2 + 1)}.$$

Hence, for $k = 2, m = 4$, the expression (36) reduces to

$$\varrho_n \approx 2n^{-1/4} \log(n)^{1/4}. \quad (38)$$

For $n = 10^7$, we obtain on average $\varrho_n \approx 0.0713$, which is significantly smaller than any of the computed $\hat{\rho}_{n,i}$. According to these considerations, we argue that our numerical computations of the covering radius are sufficiently reliable in view of (34).

Fig. 2 shows logarithmic plots of the estimated covering radii for the design points $\{P_j^i\}_{j=1}^{n_i}$, $i = 1, \dots, 14$, from Section 7.1. For comparison it also displays estimated covering radii for independent random points. We observe the desired relationship $\rho_i \asymp n_i^{-1/4}$, cf. Theorem 3.2. The estimate (38) of the expected covering radius for random points becomes more accurate for large n , and our numerical experiments support our theoretical results.

Acknowledgment

The authors have been funded by the Vienna Science and Technology Fund (WWTF) through project VRG12-009.

References

- [1] P.-A. Absil, R. Mahony, R. Sepulchre, *Optimization Algorithms on Matrix Manifolds*, Princeton University Press, 2008.
- [2] C. Bachoc, Designs, groups and lattices, *J. Theor. Nombres Bordeaux* (2005) 25–44.
- [3] C. Bachoc, E. Bannai, R. Coulangéon, Codes and designs in Grassmannian spaces, *Discrete Math.* 277 (2004) 15–28.
- [4] C. Bachoc, R. Coulangéon, G. Nebe, Designs in Grassmannian spaces and lattices, *J. Algebraic Combin.* 16 (2002) 5–19.
- [5] C. Bachoc, M. Ehler, Tight p -fusion frames, *Appl. Comput. Harmon. Anal.* 35 (1) (2013) 1–15.
- [6] A. Bondarenko, D. Radchenko, M. Viazovska, Optimal asymptotic bounds for spherical designs, *Ann. of Math.* 178 (2) (2013) 443–452.
- [7] L. Bandolini, C. Choirat, L. Colzani, G. Gigante, R. Seri, G. Travaglini, Quadrature rules and distribution of points on manifolds, *Ann. Sc. Norm. Super. Pisa Cl. Sci. XIII* (4) (2014) 889–923.
- [8] J.S. Brauchart, J. Dick, E.B. Saff, I.H. Sloan, Y.G. Wang, R.S. Womersley, Covering of spheres by spherical caps and worst-case error for equal weight cubature in Sobolev spaces, *J. Math. Anal. Appl.* 431 (2) (2015) 782–811.
- [9] J.S. Brauchart, P.J. Grabner, Distributing many points on spheres: Minimal energy and designs, *J. Complexity* 31 (2015) 293–326.
- [10] J.S. Brauchart, A.B. Reznikov, E.B. Saff, I.H. Sloan, Y.G. Wang, R.S. Womersley, Random point sets on the sphere - hole radii, covering, and separation, *Exp. Math.* (2016) 1–20.
- [11] J.S. Brauchart, E.B. Saff, I.H. Sloan, R.S. Womersley, QMC designs: Optimal order quasi Monte Carlo integration schemes on the sphere, *Math. Comp.* 83 (2014) 2821–2851.
- [12] A. Breger, M. Ehler, M. Gräf, Quasi Monte Carlo integration and kernel-based function approximation on Grassmannians, in: *Frames and Other Bases in Abstract and Function Spaces: Novel Methods in Harmonic Analysis*, vol. 1, Birkhäuser/Springer, 2017.
- [13] Y. Chikuse, Statistics on special manifolds, in: *Lecture Notes in Statistics*, Springer, New York, 2003.
- [14] C. Choirat, R. Seri, Numerical properties of generalized discrepancies on spheres of arbitrary dimension, *J. Complexity* 29 (2013) 216–235.
- [15] S.B. Damelin, P.J. Grabner, Energy functionals numerical integration and asymptotic equidistribution on the sphere, *J. Complexity* 19 (2003) 231–246.
- [16] A.W. Davis, Spherical functions on the Grassmann manifold and generalized Jacobi polynomials - Part 2, *Linear Algebra Appl.* 289 (1–3) (1999) 95–119.
- [17] P. de la Harpe, C. Pache, Cubature formulas, geometrical designs, reproducing kernels, and Markov operators, in: *Infinite Groups: Geometric, Combinatorial and Dynamical Aspects* (Basel), vol. 248, Birkhäuser, 2005, pp. 219–267.
- [18] D. Dung, T. Ullrich, Lower bounds for the integration error for multivariate functions with mixed smoothness and optimal Fibonacci cubature for functions on the square, *Math. Nachr.* 288 (7) (2015) 743–762.
- [19] M. Ehler, M. Gräf, Reproducing kernels for the irreducible components of polynomial spaces on unions of Grassmannians, *arXiv* (2017).
- [20] U. Etayo, J. Marzo, J. Ortega-Cerdà, Asymptotically optimal designs on compact algebraic manifolds, *J. Monatsh. Math.* (2018).
- [21] F. Filbir, H.N. Mhaskar, A quadrature formula for diffusion polynomials corresponding to a generalized heat kernel, *J. Fourier Anal. Appl.* 16 (5) (2010) 629–657.
- [22] F. Filbir, H.N. Mhaskar, Marcinkiewicz–Zygmund measures on manifolds, *J. Complexity* 27 (6) (2011) 568–596.
- [23] M. Gräf, *Efficient Algorithms for the Computation of Optimal Quadrature Points on Riemannian Manifolds*, Universitätsverlag Chemnitz, 2013.
- [24] P. Hellekalek, P. Kritzer, F. Pillichshammer, Open type quasi-Monte Carlo integration based on Halton sequences in weighted Sobolev spaces, *J. Complexity* 33 (2016) 169–189.
- [25] A. Hinrichs, L. Markhasin, J. Oettershagen, T. Ullrich, Optimal quasi-Monte Carlo rules on order 2 digital nets for the numerical integration of multivariate periodic functions, *Numer. Math.* (2015).
- [26] A. Hinrichs, E. Novak, M. Ullrich, H. Wozniakowski, The curse of dimensionality for numerical integration of smooth functions II, *J. Complexity* 30 (2) (2014) 117–143.
- [27] L. Hörmander, *The Analysis of Linear Partial Differential Operators*, I, II, III, IV, Springer Verlag, 1983 1985.
- [28] D. Krieg, E. Novak, A universal algorithm for multivariate integration, *Found. Comput. Math.* (2016).
- [29] H.N. Mhaskar, Approximate quadrature measures on data-defined spaces, 2017. *arXiv:1612.02368*.

- [30] H. Niederreiter, *Random Number Generation and Quasi-Monte Carlo Methods*, Society for Industrial and Applied Mathematics, Philadelphia, PA, USA, 1992.
- [31] H. Niederreiter, Some current issues in quasi-Monte Carlo methods, *J. Complexity* 19 (3) (2003) 428–433.
- [32] E. Novak, H. Wozniakowski, *Tractability of Multivariate Problems. Volume II*, in: EMS Tracts in Mathematics, vol 12, EMS Publishing House, Zürich, 2010.
- [33] A. Reznikov, E.B. Saff, The covering radius of randomly distributed points on a manifold, *Int. Math. Res. Not.* rnv342 (2015).
- [34] P. Seymour, T. Zaslavsky, Averaging sets: A generalization of mean values and spherical designs, *Adv. Math.* 52 (1984) 213–240.
- [35] I.H. Sloan, R.S. Womersley, Extremal systems of points and numerical integration on the sphere, *Adv. Comput. Math.* 21 (2004) 107–125.
- [36] H. Triebel, *Theory of Function Spaces II*, Birkhäuser, Basel, 1992.
- [37] T. Ullrich, Optimal cubature in Besov spaces with dominating mixed smoothness on the unit square, *J. Complexity* 30 (2014) 72–94.

(P2)

On Orthogonal Projections for Dimension Reduction and Applications in Augmented Target Loss Functions for Learning Problems

Journal: Journal of Mathematical Imaging and Vision (JMIV), Springer

Authors: A. Breger, J. I. Orlando, P. Harar, M. Doerfler, S. Klimscha, C. Grechenig, B. S. Gerendas, U. Schmidt-Erfurth, and M. Ehler

Published: 2019

Links: <http://link.springer.com/article/10.1007/s10851-019-00902-2>
arXiv: <https://arxiv.org/abs/1901.07598>

DOI: 10.1007/s10851-019-00902-2

Contribution: Conceptual idea; Theory development and proofs; Design and performance of numerical experiments; Main writing.



On Orthogonal Projections for Dimension Reduction and Applications in Augmented Target Loss Functions for Learning Problems

A. Breger¹ · J. I. Orlando² · P. Harar³ · M. Dörfler¹ · S. Klimescha² · C. Grechenig² · B. S. Gerendas^{1,2} · U. Schmidt-Erfurth² · M. Ehler¹

Received: 30 December 2018 / Accepted: 16 August 2019
© The Author(s) 2019

Abstract

The use of orthogonal projections on high-dimensional input and target data in learning frameworks is studied. First, we investigate the relations between two standard objectives in dimension reduction, preservation of variance and of pairwise relative distances. Investigations of their asymptotic correlation as well as numerical experiments show that a projection does usually not satisfy both objectives at once. In a standard classification problem, we determine projections on the input data that balance the objectives and compare subsequent results. Next, we extend our application of orthogonal projections to deep learning tasks and introduce a general framework of augmented target loss functions. These loss functions integrate additional information via transformations and projections of the target data. In two supervised learning problems, clinical image segmentation and music information classification, the application of our proposed augmented target loss functions increases the accuracy.

Keywords Orthogonal Projection · Dimension reduction · Preservation of data characteristics · Supervised learning · Target features

1 Introduction

Linear dimension reduction is commonly used for preprocessing of high-dimensional data in complicated learning frameworks to compress and weight important data features. In contrast to nonlinear approaches, the use of orthogonal projections is computationally cheap, since it corresponds to a simple matrix multiplication. Conventional approaches apply specific projections that preserve essential information and complexity within a more compact representation. The projector is usually selected by optimizing distinct objectives, such as information preservation of the sample variance or of pairwise relative distances. Widely used orthogonal projections for dimension reduction are variants of the prin-

cipal component analysis (PCA) that maximize the variance of the projected data [37]. Preservation of relative pairwise distances asks for a near-isometric embedding, and random projections guarantee this embeddings with high probability, cf. [5,15] and see also [1,6,12,27,30,35]. The use of random projections is especially favorable for large, high-dimensional data [48], since the computational complexity is just $O(dkm)$, e.g., using the construction in [1], with $d, k \in \mathbb{N}$ being the original and lower dimensions and $m \in \mathbb{N}$ the number of samples. In contrast, PCA needs $O(d^2m) + O(d^3)$ operations [24]. Moreover, tasks that do not have all data available at once, e.g., data streaming, ask for dimension reduction methods that are independent of the data.

In the present manuscript, we study orthogonal projections regarding the interplay between

- (O1) preservation of variance,
- (O2) preservation of pairwise relative distances,

aiming for a sufficient lower-dimensional data representation. We shall consider the Euclidean distance exclusively since it is most widely used in applications, especially

A. Breger
anna.breger@univie.ac.at

¹ Department of Mathematics, University of Vienna, Vienna, Austria

² Department of Ophthalmology, Medical University of Vienna, Vienna, Austria

³ Department of Telecommunications, Brno University of Technology, Brno, Czech Republic

for error estimation. On manifolds, the geodesic distance is locally equivalent to the Euclidean distance. The two objectives (O1) and (O2) are directly addressed by PCA (O1) and random projections (O2). We achieve the following goals: First, we clarify mathematically and numerically that the two objectives are competing, i.e., PCA and random projections preserve different kinds of information. Depending on the objectives, we discuss beneficial choices of orthogonal projections and numerically find a balancing projector for a given data set. Finally, we define a general framework of augmented target (AT) loss functions for deep neural networks that integrate information about target characteristics via features and projections. We observe that our proposed methodology can increase the accuracy in two deep learning problems.

In contrast to conventional approaches, we study the joint behavior of the two objectives with respect to the entire set of orthogonal projectors. By analyzing the correlation between the variance and pairwise relative distances of projected data, we observe that (O1) and (O2) are competing and usually cannot be reached at the same time. In classification experiments with support vector machine and shallow neural networks, we investigate heuristic choices of projections applied to the input features.

In view of learning frameworks, we utilize features and projections on target data. The class of augmented target loss functions incorporates suitable transformations and projections that provide beneficial representations of the target space. It is applied in two supervised deep learning problems dealing with real-world data.

The first experiment is a clinical image segmentation problem in optical coherence tomography (OCT) data of the human retina. Related principles of dimension reduction for other clinical classification problems in OCT have already been successfully applied in [9]. In the second experiment, we aim to categorize musical instruments based on their spectrogram; see [19] for related results. Our utilized augmented target loss functions can increase the accuracy in both experiments.

The outline is as follows. In Sect. 2, we address the analysis of the competing objectives and Theorem 2.5 yields the asymptotic correlation between variance and pairwise relative distances of projected data. Section 3 prepares for the numerical investigations by recalling t -designs as considered in [10], enabling subsequent numerics. Heuristic investigations on projected input used in a straightforward classification task are presented in Sect. 4. Our framework of augmented target loss functions as modified standard loss functions for deep learning is introduced in Sect. 5. Finally, in Sects. 6 and 7 we present classification experiments on OCT images and musical instruments using aligned augmented target loss functions.

2 Dimension Reduction with Orthogonal Projections

To reduce the dimension of a high-dimensional data set $x = \{x_i\}_{i=1}^m \subset \mathbb{R}^d$, we map x onto a lower-dimensional affine linear subspace $\bar{x} + V$, where $\bar{x} := \frac{1}{m} \sum_{i=1}^m x_i$ is the sample mean and V is a k -dimensional linear subspace of $\mathbb{R}^d$ with $k < d$. This mapping is performed by an orthogonal projector $p \in \mathcal{G}_{k,d}$, where

$$\mathcal{G}_{k,d} := \{p \in \mathbb{R}^{d \times d} : p^2 = p, p^\top = p, \text{rank}(p) = k\}$$

denotes the Grassmannian, so that the lower-dimensional data representation is

$$\{\bar{x} + p(x_i - \bar{x})\}_{i=1}^m \subset \bar{x} + V, \quad (2.1)$$

with $\text{range}(p) = V$. A suitable choice of p within $\mathcal{G}_{k,d}$ depends on further objectives, i.e., which kind of information preservation shall be favored for subsequent analysis tasks. In the following, we consider two objectives associated with popular choices of orthogonal projectors for dimension reduction, in particular, random projectors and PCA. We will first observe that the two objectives are competing, especially in high dimensions, and then discuss consequences.

2.1 Objective (O1)

The total sample variance $\text{tvar}(x)$ of $x = \{x_i\}_{i=1}^m \subset \mathbb{R}^d$ is the sum of the corrected variances along each dimension:

$$\text{tvar}(x) := \frac{1}{m-1} \sum_{i=1}^m \|x_i - \bar{x}\|^2. \quad (2.2)$$

PCA aims to construct $p \in \mathcal{G}_{k,d}$, such that the total sample variance of (2.1) is maximized among all projectors in $\mathcal{G}_{k,d}$. For other equivalent optimality criteria, we refer to [49].

The total sample variance of $px = \{px_i\}_{i=1}^m \subset V$ coincides with the one of (2.1) and satisfies

$$\text{tvar}(px) \leq \text{tvar}(x)$$

for all $p \in \mathcal{G}_{k,d}$. Thus, PCA achieves optimal variance preservation. The total variance (2.2) can also be expressed via pairwise absolute distances:

$$\text{tvar}(x) = \frac{1}{m(m-1)} \sum_{i < j} \|x_i - x_j\|^2. \quad (2.3)$$

Equally, it holds that

$$\text{tvar}(px) = \frac{1}{m(m-1)} \sum_{i < j} \|p(x_i) - p(x_j)\|^2, \quad (2.4)$$

which reveals that PCA maximizes the sample mean of the projected pairwise absolute distances.

2.2 Objective (O2)

In contrast to pairwise absolute distances, the Johnson–Lindenstrauss lemma targets the global property of preservation of pairwise relative distances:

Lemma 2.1 (Johnson–Lindenstrauss, cf. [15,35]). *For any $0 < \epsilon < 1$, any $k \leq d$, $m \in \mathbb{N}$, with*

$$\frac{4 \log(m)}{\epsilon^2/2 - \epsilon^3/3} \leq k,$$

and any set $\{x_i\}_{i=1}^m \subset \mathbb{R}^d$, there is a projector $p \in \mathcal{G}_{k,d}$ such that

$$(1-\epsilon) \|x_i - x_j\|^2 \leq \frac{d}{k} \|p(x_i) - p(x_j)\|^2 \leq (1+\epsilon) \|x_i - x_j\|^2 \quad (2.5)$$

holds for all $i < j$.

For small $\epsilon > 0$, the projector p in Lemma 2.1 yields that all of the $\frac{m(m-1)}{2}$ pairwise relative distances

$$\left\{ \frac{d}{k} \frac{\|p(x_i) - p(x_j)\|^2}{\|x_i - x_j\|^2} : i < j \right\} \quad (2.6)$$

are close to 1, i.e., the projection p preserves all scaled pairwise relative distances well. A good choice of p in Lemma 2.1 is based on random projectors¹ $P \sim \lambda_{k,d}$, where $\lambda_{k,d}$ denotes the unique orthogonally invariant probability measure on $\mathcal{G}_{k,d}$. The following theorem is essentially proved by following the lines of the proof of Lemma 2.1 in [15] after replacing the constant 4 with $(2 + \tau)2$ in the respective bound on k .

Theorem 2.2 *For any $0 < \epsilon < 1$, any $k \leq d$, $m \in \mathbb{N}$ and any $0 < \tau$ with*

$$\frac{(2 + \tau)2 \log(m)}{\epsilon^2/2 - \epsilon^3/3} \leq k,$$

and any set $\{x_i\}_{i=1}^m \subset \mathbb{R}^d$, the random projector $P \sim \lambda_{k,d}$ satisfies

$$\left\{ \frac{d}{k} \frac{\|P(x_i) - P(x_j)\|^2}{\|x_i - x_j\|^2} : i < j \right\} \in [1 - \epsilon, 1 + \epsilon] \quad (2.7)$$

with probability at least $1 - \frac{1}{m^\tau} + \frac{1}{m^{\tau+1}}$.

¹ We use lower case letters for samples and upper case letters for random vectors/matrices.

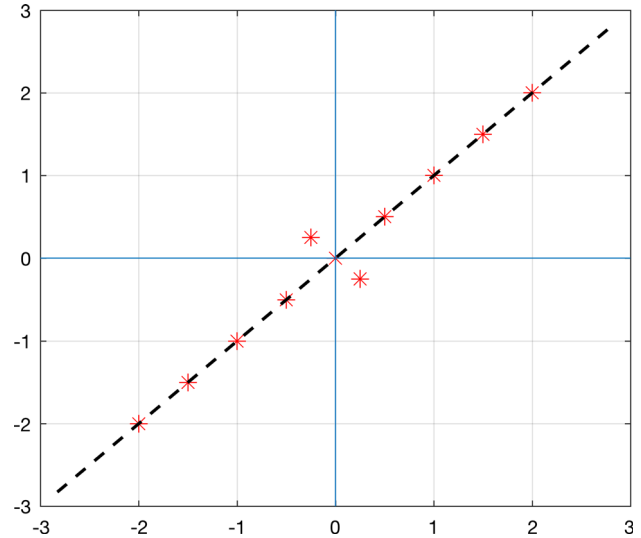


Fig. 1 A trivial example of PCA distorting smaller distances. Choosing the first principal component, PCA projects the two-dimensional data points * onto the plane of the first eigendirection (—). The Euclidean distances of the points lying on the diagonal are preserved, whereas the two points with smaller distances are projected onto a single point (the origin)

The theorem tells that the preservation property of pairwise relative distances is achieved with high probability when choosing a random projection according to $k; d$. Note that the random choice is completely independent from the actual data set.

2.3 Competing Objectives

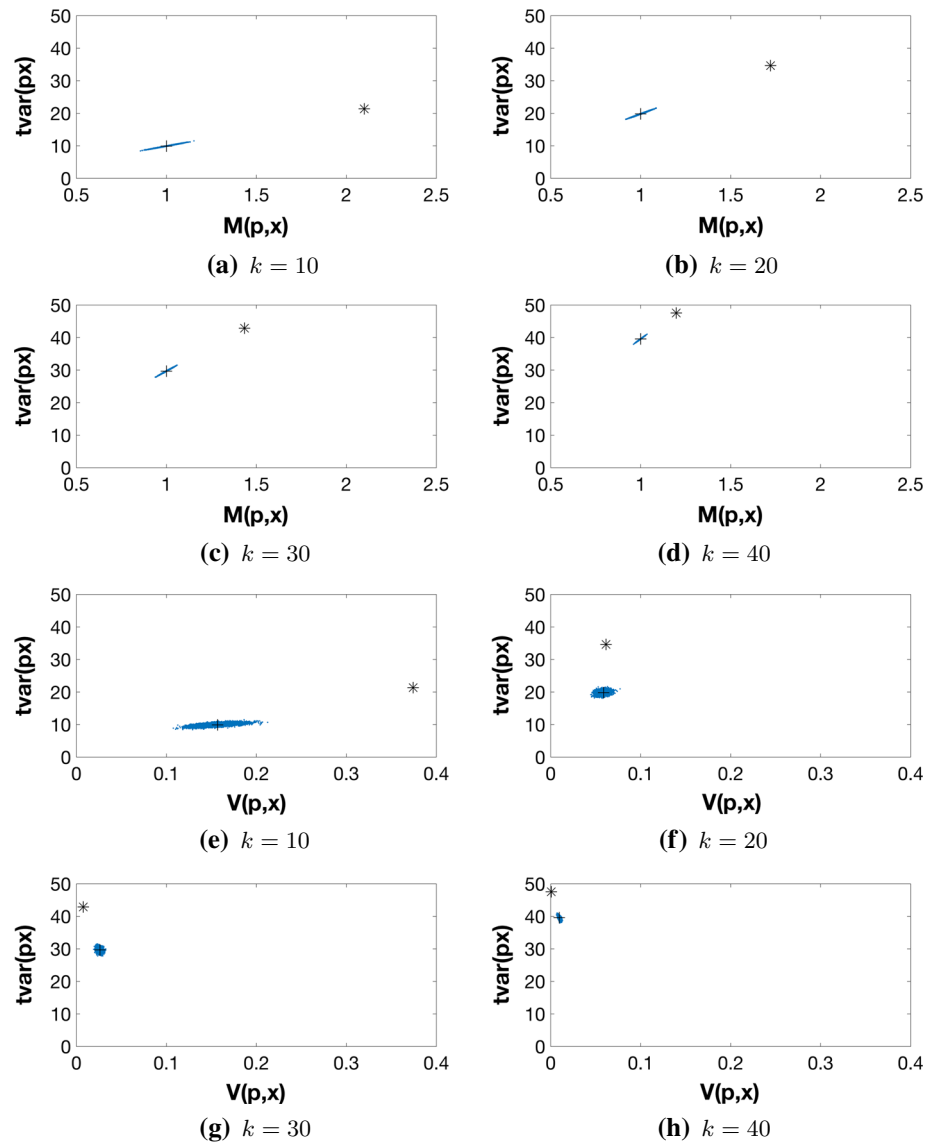
A projector p satisfying the near-isometry property (2.5) implies

$$(1 - \epsilon) \frac{k}{d} \text{tvar}(x) \leq \text{tvar}(px) \leq (1 + \epsilon) \frac{k}{d} \text{tvar}(x),$$

so that the total variance of the projected data px is not preserved for $k < d$. In particular, with high probability a random projector $P \sim \lambda_{k,d}$ does not suit the objective of maximizing the total variance, and we even observe $\mathbb{E} \text{tvar}(Px) = \frac{k}{d} \text{tvar}(x)$; see (A.2) in the Appendix. PCA does not guarantee any local geometric property, and distances between pairs of points can be arbitrarily distorted [1]; see [39] for more robust PCA. The preservation of larger distances is favored since PCA maximizes (2.4) among all $p \in \mathcal{G}_{k,d}$ and $\|p(x_i) - p(x_j)\| \leq \|x_i - x_j\|$ holds for all $i < j$. Close but distinct points could even be projected onto a single point, which violates the preservation of pairwise relative distances; see Fig. 1.

To more quantitatively understand the relation between the two competing objectives, we consider the sample mean

Fig. 2 Competing properties: 10,000 random projections $p \sim \lambda_{k,50}$ versus PCA (*), plotted concerning $\text{tvar}(px)$, $\mathcal{M}(p, x)$ and $\mathcal{V}(p, x)$. The normal distributed fixed data set x has total variance $\text{tvar}(x) = 49.5$. Random projections cluster around their expectation values (2.10), (2.11) and (2.12), marked by +



and the uncorrected sample variance of the pairwise relative distances (2.6):

$$\mathcal{M}(p, x) := \frac{2}{m(m-1)} \sum_{i < j} \frac{d}{k} \frac{\|p(x_i - x_j)\|^2}{\|x_i - x_j\|^2}, \quad (2.8)$$

$$\mathcal{V}(p, x) := \frac{2}{m(m-1)} \sum_{i < j} \frac{d^2}{k^2} \frac{\|p(x_i - x_j)\|^4}{\|x_i - x_j\|^4} - \mathcal{M}(p, x)^2. \quad (2.9)$$

Recall that good preservation of the relative pairwise distances in (2.6) asks for $\mathcal{M}(p, x)$ being close to 1 and the variance $\mathcal{V}(p, x)$ being small. In the following, we analyze $\text{tvar}(px)$, $\mathcal{M}(p, x)$ and $\mathcal{V}(p, x)$ and their expectations for random $P \sim \lambda_{k,d}$.

In Fig. 2, we see a simple numerical experiment, where we first create an independent, normally distributed fixed data

set $\{x_i\}_{i=1}^m$ with $x_i \in \mathbb{R}^d$ for $i = 1, \dots, m$ and $m = 100$, $d = 50$. We then compute PCA, for $k = 10, 20, 30, 40$, as well as $n = 10,000$ random projections p distributed according to $\lambda_{k,50}$. In Fig. 2a–d, we can see that the more the k differs from d , the more the PCA and random projections differ concerning $\text{tvar}(px)$ and $\mathcal{M}(p, x)$. Those differences may lead to diverse behavior in subsequent data analysis. Moreover, we compare $\mathcal{M}(p, x)$ and $\mathcal{V}(p, x)$ in Fig. 2e–h for the different k . We can see that again when k is much smaller than d , random projections and PCA differ more concerning the variance of pairwise distances $\mathcal{V}(p, x)$. For $k = 10$, the variance for PCA is higher in comparison with random projections (Fig. 2e); for $k = 40$ vice versa (Fig. 2h). Note that the theoretical bounds stated in Theorem 2.1 are much higher than the dimensions k used in the experiments, but the projections still preserve relative pairwise distances

very well. In [7], similar observations were made on empirical experiments with image and text data.

The amount of variance kept in the principal components comparing real-world and random data has been experimentally studied, e.g., in [29] and [46]. Both studies determine that the difference occurs mainly in the first principal component.

Remark 2.3 In the numerical example, we compare random projections and PCA directly, serving as the corresponding projections to the objectives (O1) and (O2). We observe that even for not so high-dimensional ($d = 50$) data x and $k \ll d/2$, PCA severely loses information in terms of total variance, i.e., more than 50% for $k = 10$, and more importantly loses much more information on pairwise relative distances than random projections. If both types of information are of interest, pairwise relative distances and high total variance, one should therefore favor random projections over PCA for $k \ll d/2$ to balance the two objectives (O1) and (O2) and vice versa. Note that with a large amount of data one might still want to favor random projectors since their construction is computationally much cheaper and independent from the data. On the other hand, if objective (O2) is negligible, e.g., tasks with very noisy data, then PCA would be the favorable choice for all k .

Information of data can be quantified and expressed in different ways. One crucial part in dimension reduction is the decision of what kind of information shall be kept, which depends on several parameters including the quality of the data and the analysis task. Variants of PCA, focusing on the preservation of variance, have been widely used in real-world problems with big success, especially in denoising, when the preservation of all pairwise relative distances may be counterproductive, e.g., in dMRI imaging [51] and color filter array images [56]. Drawbacks are the necessity for all data being available from the start and the high computational costs. For very high-dimensional and large data sets, the computation of PCA is often not feasible. Besides the huge benefit of data independence and low computational cost when using random projections, the near-isometry property often allows to establish that the solution found in the low-dimensional space is a good approximation to the solution in the original space [1, 34].

Algorithms in machine learning often need or benefit from sufficient estimates of pairwise distances, e.g., approximate nearest-neighbor problems, supervised classification [27] and subspace clustering [26]. In [32], algorithmic applications of near-isometry embeddings have been introduced. In [7], random projections have been successfully applied to noisy and noiseless text and image data. The experimental studies include the comparison of preservation of pairwise distances between random projections and PCA. The results coincide with our observations that for $k > d/2$ PCA is able

to preserve the pairwise distances sufficiently, whereas for $k < d/2$ PCA distorts them. The smaller the k , the worse the distortion, whereas random projections preserve similarities still well for very small k , while being computationally much cheaper than PCA. One should point out again that favoring preservation of pairwise distances relies on the accuracy of the original distances.

PCA and random projections are orthogonal projections favoring two different aims. We want to study in the context of the whole set of orthogonal projections if the two objectives (O1) and (O2) could be reached at the same time. We will see that the objectives act competing, and therefore we suggest a balancing projector for tasks that benefit from both objectives.

2.4 Covariances and Correlation Between Competing Objectives

For further mathematical analysis, we first introduce a more general class of probability measures on $\mathcal{G}_{k,d}$ that resemble $\lambda_{k,d}$ sufficiently well:

Definition 2.4 A Borel probability measure λ on $\mathcal{G}_{k,d}$ is called a *cubature measure of strength t* if

$$\int_{\mathcal{G}_{k,d}} f(p) d\lambda_{k,d}(p) = \int_{\mathcal{G}_{k,d}} f(p) d\lambda(p), \quad \text{for all } f \in \text{Pol}_t(\mathbb{R}^{d^2}),$$

where $\text{Pol}_t(\mathbb{R}^{d^2})$ denotes the set of multivariate polynomials of total degree t in d^2 variables.

Existence of cubature measures is studied, for instance, in [17]. For random P , we now determine the expectation values for our three quantities of interest: $\text{tvar}(Px)$, $\mathcal{M}(P, x)$ and $\mathcal{V}(P, x)$. If $P \sim \lambda$ and λ is a cubature measure of strength at least 2, the identities (A.2) and (A.3) in the Appendix and a short calculation yield

$$\mathbb{E} \text{tvar}(Px) = \frac{k}{d} \text{tvar}(x), \quad (2.10)$$

$$\mathbb{E} \mathcal{M}(P, x) = 1, \quad (2.11)$$

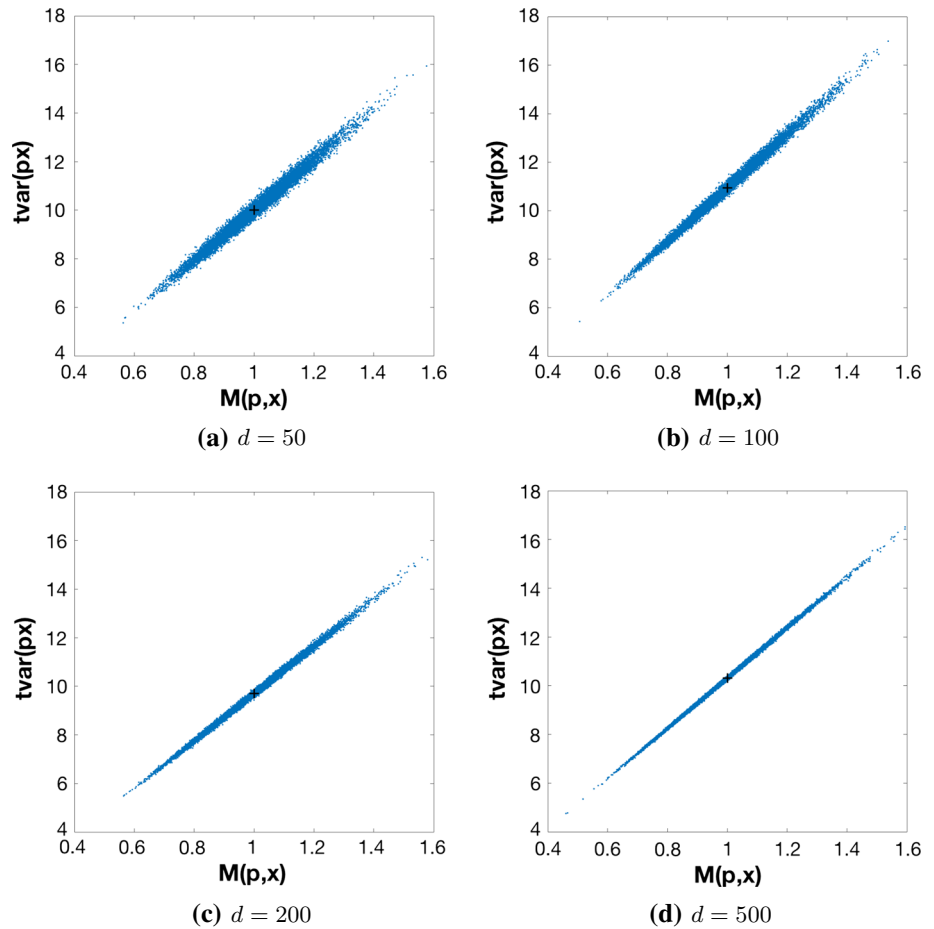
$$\mathbb{E} \mathcal{V}(P, x) = a_{k,d} \left(1 - \frac{4}{m^2(m-1)^2} \sum_{\substack{i < j \\ l < r}} \left\langle \frac{x_i - x_j}{\|x_i - x_j\|}, \frac{x_l - x_r}{\|x_l - x_r\|} \right\rangle^2 \right), \quad (2.12)$$

where $a_{k,d} = \frac{2d(d-k)}{k(d-1)(d+2)}$. The expected sample variance in (2.12) satisfies

$$\mathbb{E} \mathcal{V}(P, x) \leq a_{k,d} \longrightarrow \frac{2}{k}, \quad \text{for } d \rightarrow \infty.$$

This asymptotic bound relates to Theorem 2.2 and alludes to a near-isometry property of the type (2.7) for k sufficiently large.

Fig. 3 For $x = \{x_i\}_{i=1}^{10} \subset \mathbb{R}^d$ with independent, normal distributed entries, we independently sample 10,000 random projectors p from $\lambda_{10,d}$ and plot $\mathcal{M}(p, x)$ versus $\text{tvar}(px)$. The expectation values with respect to $P \sim \lambda$ are marked with $+$. The correlation is already 0.9916 for $d = 50$ and grows further when d increases, namely with values 0.9961, 0.9985, 0.9996 for $d = 100, 200, 500$



The following theorem provides a lower bound for random P on the population correlation

$$\text{Corr}(\mathcal{M}(P, x), \text{tvar}(Px)) = \frac{\text{Cov}(\mathcal{M}(P, x), \text{tvar}(Px))}{\sqrt{\text{Var}(\mathcal{M}(P, x))} \sqrt{\text{Var}(\text{tvar}(Px))}}. \quad (2.13)$$

Theorem 2.5 Let $x = \{x_i\}_{i=1}^m \subset \mathbb{R}^d$ be pairwise different and let $P \sim \lambda$, with λ being a cubature measure of strength at least 2. For $d \geq \frac{m(m-1)}{2}$, the correlation (2.13) is bounded from below by

$$\frac{\min_{i \neq j} \|x_i - x_j\|^2}{\max_{i \neq j} \|x_i - x_j\|^2} - \frac{m(m-1)}{2d} \cdot \frac{\max_{i \neq j} \|x_i - x_j\|^2}{\min_{i \neq j} \|x_i - x_j\|^2}. \quad (2.14)$$

If $\{x_i\}_{i=1}^m \subset \mathbb{R}^d$ are random points, whose entries are independent, identically distributed with finite 4-th moments that are uniformly bounded in d , then (2.14) converges towards 1 in probability for $d \rightarrow \infty$.

The strong correlation for large dimensions d in the second part of Theorem 2.5 suggests that increasing $\text{tvar}(Px)$ may also lead to increasing $\mathcal{M}(P, x)$; see Fig. 3 for illustration. Thus, large projected total variance $\text{tvar}(Px)$ and the preservation of scaled pairwise distances, i.e., $\mathcal{M}(P, x)$ being close

to 1, are competing properties. As discussed in Sect. 2.3, the choice of which kind of information is favorable to preserve depends on the data and the task, e.g., denoising (O1) and nearest-neighbor classification (O2). PCA and random projections are extreme in preserving either (O1) or (O2). We will heuristically study the behavior of orthogonal projections balancing both objectives in the next section and will state a numerical experiment where a balancing projector yields the highest classification accuracy.

Remark 2.6 The second part of Theorem 2.5 relates to the well-known fact that random vectors in high dimensions are almost orthogonal [4], and standard concentration of measure arguments may lead to more quantitative statements, cf. [52].

3 Preparations for Numerical Experiments

For the numerical experiments, we need finite sets of projectors that represent the overall space well, i.e., cover $\mathcal{G}_{k,d}$ properly.

3.1 Optimal Covering Sequences

Let the *covering radius* of a set $\{p_l\}_{l=1}^n \subset \mathcal{G}_{k,d}$ be denoted by

$$\varrho(\{p_l\}_{l=1}^n) := \sup_{p \in \mathcal{G}_{k,d}} \min_{1 \leq l \leq n} \|p - p_l\|_F, \quad (3.1)$$

where $\|\cdot\|_F$ is the Frobenius norm. The smaller the covering radius, the better the set $\{p_l\}_{l=1}^n$ represents the entire space $\mathcal{G}_{k,d}$, i.e., there are smaller holes and the points $\{p_l\}_{l=1}^n$ are better distributed within $\mathcal{G}_{k,d}$. Following Lemma 2.1, we can connect finite sets of projections and their covering radius to the near-isometry property:

Lemma 3.1 *Let $\{p_l\}_{l=1}^n \subset \mathcal{G}_{k,d}$ and denote $\varrho := \varrho(\{p_l\}_{l=1}^n)$. For any $0 < \epsilon < 1$, any $m, k, d \in \mathbb{N}$ with*

$$\frac{4 \log(m)}{\epsilon^2/2 - \epsilon^3/3} \leq k \leq d,$$

and any $\{x_i\}_{i=1}^m \subset \mathbb{R}^d$, there is $l_0 \in \{1, \dots, n\}$ such that

$$(1 - \delta) \|x_i - x_j\|^2 \leq \frac{d}{k} \|p_{l_0}(x_i) - p_{l_0}(x_j)\|^2 \leq (1 + \delta) \|x_i - x_j\|^2, \quad i < j, \quad (3.2)$$

where $\delta = \epsilon + 2\varrho\sqrt{\frac{(1+\epsilon)d}{k}} + \frac{d}{k}\varrho^2$.

Proof Given an arbitrary projector $p \in \mathcal{G}_{k,d}$, there is an index $l_0 \in \{1, \dots, n\}$ such that

$$\|p_{l_0}x - px\| \leq \|p_{l_0} - p\|_F \|x\| \leq \varrho \|x\|, \quad x \in \mathbb{R}^d.$$

From here, standard computations imply Lemma 3.1. We omit the details. $\square$

The accuracy of the near-isometry property in (3.2) depends on the covering radius. Therefore, a set $\{p_l\}_{l=1}^n \in \mathcal{G}_{k,d}$ with a small covering radius ϱ is more likely to contain a projector with better preservation of pairwise relative distances. According to [11], it holds that² $\varrho \gtrsim n^{-\frac{1}{k(d-k)}}$ and we shall see next how to achieve this lower bound.

A set of projectors $\{p_l\}_{l=1}^n \subset \mathcal{G}_{k,d}$ is called a *t-design* if the associated normalized atomic measure $\frac{1}{n} \sum_{l=1}^n \delta_{p_l}$ is a cubature measure of strength t (see Definition 2.4); see [44] for general existence results. Any sequence of t_i -designs $\{p_l^i\}_{l=1}^{n_i} \subset \mathcal{G}_{k,d}$ with $t_i \rightarrow \infty$ satisfies

$$\varrho_i \asymp t_i^{-1}, \quad (3.3)$$

² We use the symbols $\lesssim$ and $\gtrsim$ to indicate that the corresponding inequalities hold up to a positive constant factor on the respective right-hand side. The notation $\asymp$ means that both relations $\lesssim$ and $\gtrsim$ hold.

and moreover, the bound $n_i \gtrsim t_i^{k(d-k)}$ holds, cf. [11, 17]. To relate n_i to ϱ_i via t_i , a sequence of t_i -designs $\{p_l^i\}_{l=1}^{n_i} \subset \mathcal{G}_{k,d}$ is called a *low-cardinality design sequence* if $t_i \rightarrow \infty$ and

$$n_i \asymp t_i^{k(d-k)}, \quad i = 1, 2, \dots \quad (3.4)$$

For their existence and numerical constructions, we refer to [21] and [10, 11]. According to [11], see also (3.3) and (3.4), any low-cardinality design sequence $\{p_l^i\}_{l=1}^{n_i}$ covers asymptotically optimal, i.e.,

$$\varrho_i \asymp n_i^{-\frac{1}{k(d-k)}}.$$

Benefiting from the covering property, we will use low-cardinality design sequences as a representation of the overall space of orthogonal projectors $\mathcal{G}_{k,d}$.

3.2 Linear Least Squares Fit

With the linear least squares fit, we can directly gain information about the relation between $\mathcal{M}(p, x)$ and $\text{tvar}(px)$ for a given data set $x = \{x_i\}_{i=1}^m \subset \mathbb{R}^d$ when p varies. Given the two samples

$$\{\text{tvar}(p_1x), \dots, \text{tvar}(p_nx)\}, \quad \{\mathcal{M}(p_1, x), \dots, \mathcal{M}(p_n, x)\}, \quad (3.5)$$

the linear least squares fitting provides the best fitting straight line,

$$\text{tvar}(p_lx) \approx s \cdot \mathcal{M}(p_l, x) + \gamma, \quad l = 1, \dots, n,$$

where s and γ are determined by the sample variances and the sample covariance. If $\{p_l\}_{l=1}^n$ is a 2-design, then the sample (co)variances coincide with the respective population (co)variances for $P \sim \lambda_{k,d}$; see Appendix A.3 for further details. It follows that

$$s = \frac{\text{Cov}(\mathcal{M}(P, x), \text{tvar}(Px))}{\text{Var}(\mathcal{M}(P, x))} \quad \text{with } P \sim \lambda_{k,d}, \quad (3.6)$$

$$\gamma = \frac{k}{d} \text{tvar}(x) - s. \quad (3.7)$$

The quantities s and γ can be directly computed, where $\text{tvar}(x)$ is given by (2.2) and the covariances are stated in Corollary A.1. Note that (3.6) and (3.7) are now independent of the particular choice of $\{p_l\}_{l=1}^n$.

The correlation between the two samples (3.5) yields additional information about their relation. As before, if $\{p_l\}_{l=1}^n$ is a 2-design, then the sample correlation coincides with the population correlation (2.13) for $P \sim \lambda_{k,d}$, cf. Appendix A.3. High correlation for a specific data set x suggests that random projections and PCA preserve competing properties,

whose benefits need to be assessed for the specific subsequent task.

4 Numerical Experiments in Pattern Recognition

We investigate the impact on classification accuracy when applying specific orthogonal projections to input data. The chosen real-world data yields a straightforward classification task, serving as a toy example for comparing the accuracy of several projected input data in simple learning frameworks. Projectors are chosen from a t -design in view of $\text{tvar}(px)$ and $\mathcal{M}(p, x)$. For all computations made in this section, the ‘Neural Network’ and ‘Statistics and Machine Learning’ toolboxes in MATLAB R2017a are used.

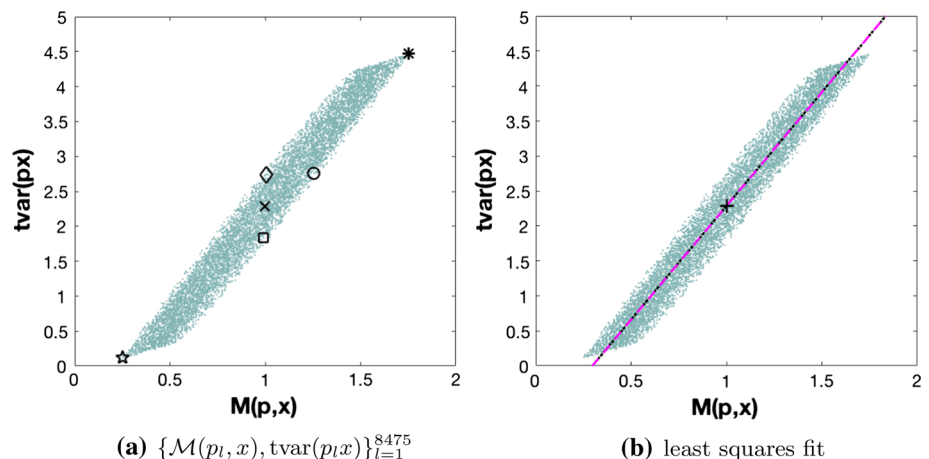
We use the publicly available *iris* data set from the UCI Repository of Machine Learning Database suitable for supervised classification learning. It consists of three classes with 50 instances each, where each class refers to a type of iris plant. The instances are described by four features resulting in the input samples $\{x_i\}_{i=1}^{150} \subset \mathbb{R}^4$ and target samples $\{y_i\}_{i=1}^{150} \subset \{0, 1\}^3$. For comparison, we classify the diverse input data with support vector machine (SVM) and three-layer neural networks (NN) with 5 and 10 hidden units (HU).

4.1 Choice of Orthogonal Projection

In the experiment, we use projections $p \in \mathcal{G}_{2,4}$ reducing the original dimension from $d = 4$ to $k = 2$. As a finite representation of the overall space, we use a t -design of strength 14 from a low-cardinality sequence (see Sect. 3.1) consisting of 8475 orthogonal projectors. Note that the dimension reduction in practice takes place by applying $q \in \mathcal{V}_{k,d}$ with $q^\top q = p \in \mathcal{G}_{k,d}$, where

$$\mathcal{V}_{k,d} := \{q \in \mathbb{R}^{k \times d} : qq^\top = I_k\}$$

Fig. 4 Projections $\{p_l\}_{l=1}^{8475} \subset \mathcal{G}_{2,4}$ from a t -design of strength 14 evaluated on the iris data set $x \subset \mathbb{R}^{4 \times 150}$



denotes the Stiefel manifold. When taking norms, p and q are interchangeable, i.e., $\|q(x)\|^2 = \|p(x)\|^2$, for all $x \in \mathbb{R}^d$. Therefore, we can use w.l.o.g. the theory developed for p .

The projections are chosen in a deterministic manner viewing the previously described competing properties. In Fig. 4, the quantities $\text{tvar}(px)$ and $\mathcal{M}(p, x)$ are pairwise plotted for all projectors in $\{p_l\}_{l=1}^{8475}$. For comparison, we choose the following projections $p \in \{p_l\}_{l=1}^{8475} \subset \mathcal{G}_{2,4}$; see Fig. 4a for a visualization.

- $p_\times$ closest to the expected values 1 and $\frac{k}{d} \text{tvar}(x)$ (see (2.10) and (2.11)),
- $p_\diamond$ preserving $\mathcal{M}(p, x) \approx 1$ and maximizing $\text{tvar}(px)$,
- $p_\square$ preserving $\mathcal{M}(p, x) \approx 1$ and minimizing $\text{tvar}(px)$,
- $p_\circ$ $\text{tvar}(px) \approx \text{tvar}(p_\diamond x)$ and maximizing $\mathcal{M}(p, x)$,
- $p_\star$ minimal $\text{tvar}(px)$,
- p_* maximal $\text{tvar}(px)$ (PCA).

4.2 Results

In Fig. 4b, we see the linear least squares fitting line, computed directly and via the slope and intercept as stated in (3.6) and (3.7). The correlation coefficient (2.13) is 0.98, which suggests that preserving the two properties is highly competing and needs to be balanced.

In Table 1, the classification results of the iris data are presented. We can see that in this comparison the projector $p_\diamond$, which corresponds to preserving $\mathcal{M}(p, x) \approx 1$ and maximizing $\text{tvar}(px)$, yields the highest and most robust results. It even yields better results than working with the original input data. The projections that preserve $\mathcal{M}(p, x) \approx 1$ but do not take care of the magnitude of the total variance yield much worse results. On the other hand, the projections that just focus on high total variance still do not yield as high results as the projection $p_\diamond$ that balances both properties.

Remark 4.1 Given a data set x , the projector $p_\diamond$ is a good choice to balance both objectives (O1) and (O2). It can be computed by directly analyzing $\{\text{tvar}(p_1 x), \dots, \text{tvar}(p_n x)\}$

Table 1 Classification results of iris data, when using projected input data in support vector machine (SVM) and shallow neural networks (NN)

Input/method	NN (10 HU)	NN (5 HU)	SVM
$\mathbf{x}$	(97.6, 1.25)	(97.5, 2.29)	(96.7, 0.15)
$\mathbf{p}_{\diamond}\mathbf{x}$	(98.4, 0.42)	(98.3, 2.15)	(97.3, 0.06)
$\mathbf{p}_{\times}\mathbf{x}$	(88, 1.73)	(87.9, 1.92)	(87.6, 0.63)
$\mathbf{p}_{\square}\mathbf{x}$	(87.3, 9.56)	(86.9, 10.81)	(87.7, 0.42)
$\mathbf{p}_{\circ}\mathbf{x}$	(96.8, 2.74)	(96.7, 1.77)	(96, 0.17)
$\mathbf{p}_{*}\mathbf{x}$ (PCA)	(96.9, 1.36)	(96.5, 4.07)	(96, 0.37)
$\mathbf{p}_{\star}\mathbf{x}$	(62.1, 44.30)	(58.9, 70.78)	(56, 0.61)

Mean and variance ($\times 10^{-4}$) of 1000 independent NN runs and 100 independent runs with tenfold cross-validation in SVM

Bold values indicate corresponding to the highest (= best) classification accuracy

and $\{\mathcal{M}(p_1, x), \dots, \mathcal{M}(p_n, x)\}$ of a finite covering $\{p_l\}_{l=1}^n$ of $\mathcal{G}_{k,d}$. For higher dimensions, an accurate representation of $\mathcal{G}_{k,d}$, in order to heuristically select $p_{\diamond}$, requires large computational costs. The least squares regression line for a 2-design, as stated in 3.7, can be directly computed with low computational cost. This offers helpful information about the interplay between (O1) and (O2).

5 Augmented Target Loss Functions

In the previous section, projectors were applied to input features of shallow neural networks. In more complex architectures, such as deep neural networks, the adaption of weights can be viewed as optimization of input features, e.g., arising features can be used for transfer learning [54]. Whereas the input data are processed and optimized in each iteration, the target data stay usually unchanged during the whole learning process, serving as a measure of accuracy. The representation of the target data is one key property for successful approximation with neural networks. Here, we will introduce a general class of loss functions, i.e., augmented target (AT) loss functions, that use projections and features to yield beneficial representations of the target space, emphasizing important characteristics.

In optimization problems, additional penalty terms are used for regularization or to enforce other constraints. In deep learning, weight decay (i.e., Tikhonov regularization) is a standard adaption of the loss function to that effect. Incorporating additional underlying information via features of the output/target data has been studied in diverse settings tailored to particular imaging applications. Perceptual loss functions have been used in [31] for image super-resolution, incorporating the comparison of high-level image features that arise from pretrained convolutional neural networks, i.e., the VGG network [45]. Deep perceptual similarity metrics

have been proposed in [20] for generating images, comparing image features instead of the original images. In [28], a similar approach was successfully used for style transfer and super-resolution, adding a network that defines loss functions. Anatomically constrained neural networks (ACNN) have been introduced in [40] and applied to cardiac image enhancement and segmentation. Their loss functions incorporate structural information by using autoencoders to gain features about lower-dimensional parametrization of the segmentation. Brain segmentation was studied in [22], where information about the desired structure has been added in the loss function via an adjacency matrix. It was used for fine-tuning the supervised learned network with unlabeled data, reducing the number of abnormalities in the segmentation.

The information of certain target characteristics can be very powerful and even replace the need of annotations in some tasks. In [47], label-free learning is approached by using just structural information of the desired output in the loss function instead of annotated target values.

In the following, we will define a general framework of loss functions that add information of target characteristics via features and projections in supervised learning tasks.

5.1 General Framework

Let the training data be input vectors $\{x_i\}_{i=1}^m \subset \mathbb{R}^r$ with associated target values $\{y_i\}_{i=1}^m \subset \mathbb{R}^s$. We consider training a neural network

$$f_{\theta} : \mathbb{R}^r \rightarrow \mathbb{R}^s,$$

where $\theta \in \mathbb{R}^N$ corresponds to the vector of all free parameters of a fixed architecture. In each optimization step for θ , the network's output $\{\hat{y}_i = f_{\theta}(x_i)\}_{i=1}^m \subset \mathbb{R}^s$ is compared with the targets $\{y_i\}_{i=1}^m$ via an underlying loss function L .

In contrast to ordinary learning problems with highly accurate target data, complicated learning tasks arising in many real-world problems do not yield sufficient results when optimizing neural networks with standard loss functions L , such as the widely used mean least squares error

$$L_{\text{MSE}}(\{y_i\}_{i=1}^m, \{\hat{y}_i\}_{i=1}^m) := \frac{1}{m} \sum_{i=1}^m \|y_i - \hat{y}_i\|^2. \quad (5.1)$$

The training data may include important information that is obvious for humans, but poorly represented within the original target data and therefore lacks consideration in the learning process. To overcome this issue, we propose to add information tailored to the particular learning problem represented by additional features of the outputs and targets.

First, we select transformations

$$T_j : \mathbb{R}^s \rightarrow \mathbb{R}^t, \quad j = 1, \dots, d,$$

to enable error estimation in transformed output/target spaces. Note that the transformations T_j are not required to be linear. However, they should be piecewise differentiable to enable subsequent optimization of the loss function with gradient methods. We shall allow for additional weighting of the transformations $T_1, \dots, T_d$ to facilitate the selection of features for a specific learning problem. The previous sections suggest that orthogonal projections can provide favorable feature combinations, which essentially turns into a weighting procedure.

To enable suitable projections, we stack the d output/target features

$$T(y_i) := \begin{pmatrix} T_1(y_i)^\top \\ \vdots \\ T_d(y_i)^\top \end{pmatrix} \in \mathbb{R}^{d \times t},$$

so that applying a projector $p \in \mathcal{G}_{k,d}$ to each column of $T(y_i)$ yields $p(T(y_i)) \in \mathbb{R}^{d \times t}$. We now define the augmented target loss function with projections by

$$L_p(\{y_i\}, \{\hat{y}_i\}) := L(\{y_i\}, \{\hat{y}_i\}) + \alpha \cdot \tilde{L}(\{p(T(y_i))\}, \{p(T(\hat{y}_i))\}), \quad (5.2)$$

where $\alpha > 0$ and L and $\tilde{L}$ correspond to conventional loss functions. Apparently, L_p depends on the choice of $p \in \mathcal{G}_{k,d}$. The projection $p(T(y_i))$ weighs the previously chosen feature transformations $T(y_i)$. Standard choices of L and $\tilde{L}$ are L_{MSE} , in which case L_p becomes

$$L_p(\{y_i\}, \{\hat{y}_i\}) = \frac{1}{m} \sum_{i=1}^m \|y_i - \hat{y}_i\|^2 + \alpha \cdot \frac{1}{m} \sum_{i=1}^m \|p(T(y_i)) - p(T(\hat{y}_i))\|_{\mathbb{F}}^2. \quad (5.3)$$

Remark 5.1 For $k = d$, the projector p is the identity. In this case, the transformations can map onto different spaces, i.e.,

$$T_j : \mathbb{R}^s \rightarrow \mathbb{R}^{t_j}, \quad j = 1, \dots, d,$$

and we can now write the standard augmented target loss function by

$$L_{\text{AT}}(\{y_i\}, \{\hat{y}_i\}) = \sum_{j=1}^d \alpha_j \cdot L^j(\{T_j(y_i)\}, \{T_j(\hat{y}_i)\}), \quad (5.4)$$

where T_1 corresponds to the identity function, $L^1, \dots, L^d$ are common loss functions and $\alpha_1, \dots, \alpha_d > 0$ are weighting parameters.

It should be mentioned that α resembles a regularization parameter. The actual minimization of (5.1) among θ is usually performed through Tikhonov-type regularization

in many standard deep neural network implementations. The formulation (5.2) adds one further variational step for beneficial output data representation.

Remark 5.2 Our proposed structure with target feature maps $T_1, \dots, T_d$ as in (5.4) relates to multitask learning, which has been successfully used in deep neural networks [13]. It handles multiple learning problems with different outputs at the same time. In contrast to multitask learning, we aim to solve a single problem but also penalize the error in transformed spaces enhancing certain target characteristics.

For the projected feature transformations in the augmented target loss function, it is not possible to identify a balancing projection p heuristically (such as $p_\diamond$ in Sect. 4), because the output y changes in each iteration when the loss function is called. In the following clinical numerical experiment we overcome this issue by choosing random projections in each optimization step and compare it to prior deterministic choices of projections, including PCA.

6 Application to Clinical Image Data

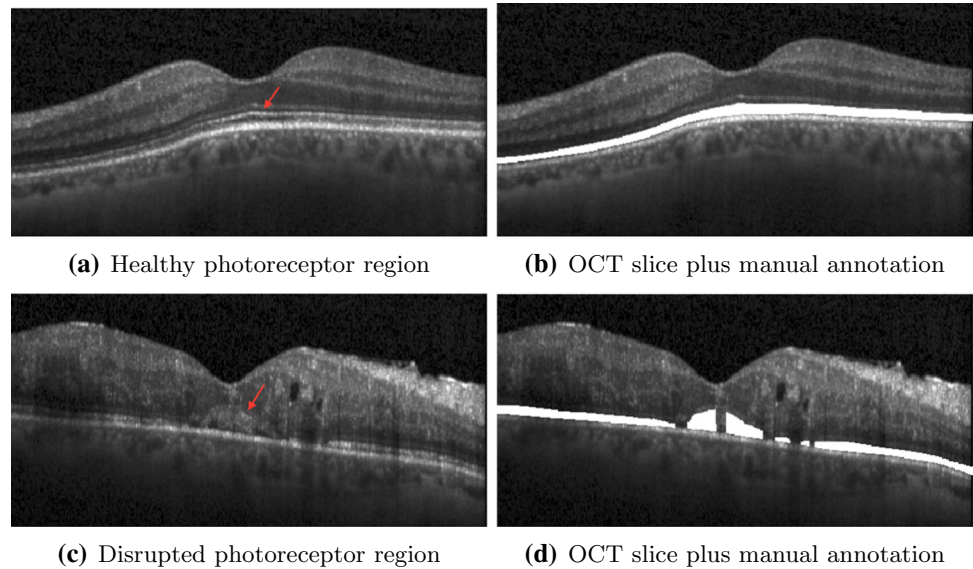
The first experiment is a clinical problem in retinal image analysis of the human eye, where the disruptions of the so-called photoreceptor layers need to be quantified in optical coherence tomography (OCT) images. The photoreceptors have been identified as the most important retinal biomarker for prediction of vision from OCT in various clinical publications, see e.g., [23]. As OCT technology advances, clinicians are not able to look at each slice of OCT themselves. (In mean, they get 250 slices per patient and have 3–5 minutes/patients including their clinical examination.) Therefore, automated classification of, for example, photoreceptor status is necessary for clinical guidance.

6.1 Data and Objective

In this application, OCT images of different retinal diseases (diabetic macular edema and retinal vein occlusion) were provided by the Vienna Reading Center recorded with the Spectralis OCT device (Heidelberg Engineering, Heidelberg, Germany). Each patient's OCT volume consists of 49 cross sections/slices (496×512 pixels) recorded in an area of 6×6 mm in the center of the human retina, which is the part of the retina responsible for vision. Each of the slices was manually annotated by a trained grader of the reading center. This is a challenging and time-consuming procedure that is not feasible in clinical routine but only in a research setting. The binary pixelwise annotations serve as target values, enabling a supervised learning framework (Fig. 5).

The objective is to accurately detect the photoreceptor layers and their disruptions pixelwise in each OCT slice by training a deep convolutional neural network with a suitable

Fig. 5 OCT provides cross-sectional visualization of the human retina



loss function. The learning problem is complicated by potentially inaccurate target annotations, as studies have shown that inconsistencies between trained graders are common, cf. [50]. Moreover, the learning task is unbalanced in the sense that there are many more slices showing none or very little disruptions. We shall observe that optimization with respect to standard loss functions performs poorly in regards to detecting disruptions. The augmented target loss function proposed in the previous section can enhance the detection.

6.2 Convolutional Neural Network Learning

We implemented our experiments using Python 3.6 with Pytorch 1.0.0. A deep convolutional neural network f_θ is trained by applying the U-Net architecture reported in [43] with a sigmoid activation function and Tikhonov regularization. A set of 20 OCT volumes (980 slices) from different patients with corresponding annotations are used for training, where four volumes were used for calibration (validation set). Another two independent test volumes were identified for evaluating the results, one without any disruptions in the photoreceptor layers, whereas the other one includes a high number of disruptions.

Each OCT slice is represented by a vector $x_i \in \mathbb{R}^r$ with $r = 496 \cdot 512$. The collection $\{x_i\}_{i=1}^m$ corresponds to all slices from the training volumes, i.e., $m = 20 \cdot 49$. Further matching the notation of the previous section, we have $r = s$ and $f_\theta : \mathbb{R}^r \rightarrow \mathbb{R}^r$ with binary target vectors $y_i \in \{0, 1\}^r$. We observe that disruptions are not identified reliably when using the least squared loss function (5.1). To overcome this issues, we use the proposed augmented target loss function with least squared losses as stated in (5.3).

To enhance disruptions within the output/target space, we heuristically choose $d = 4$ local features of the original rep-

Table 2 Comparison of AUC values for photoreceptors segmentation and disruption detection

Loss function	Photoreceptors	Disruptions
L_{MSE}	0.9720	0.4399
L_p		
$p = I_4$	0.9736	0.4686
$p_{\lambda_{2,4}}$	0.9746	0.4720
p_{PCA}	0.9716	0.5331
p_{12}	0.9755	0.5558

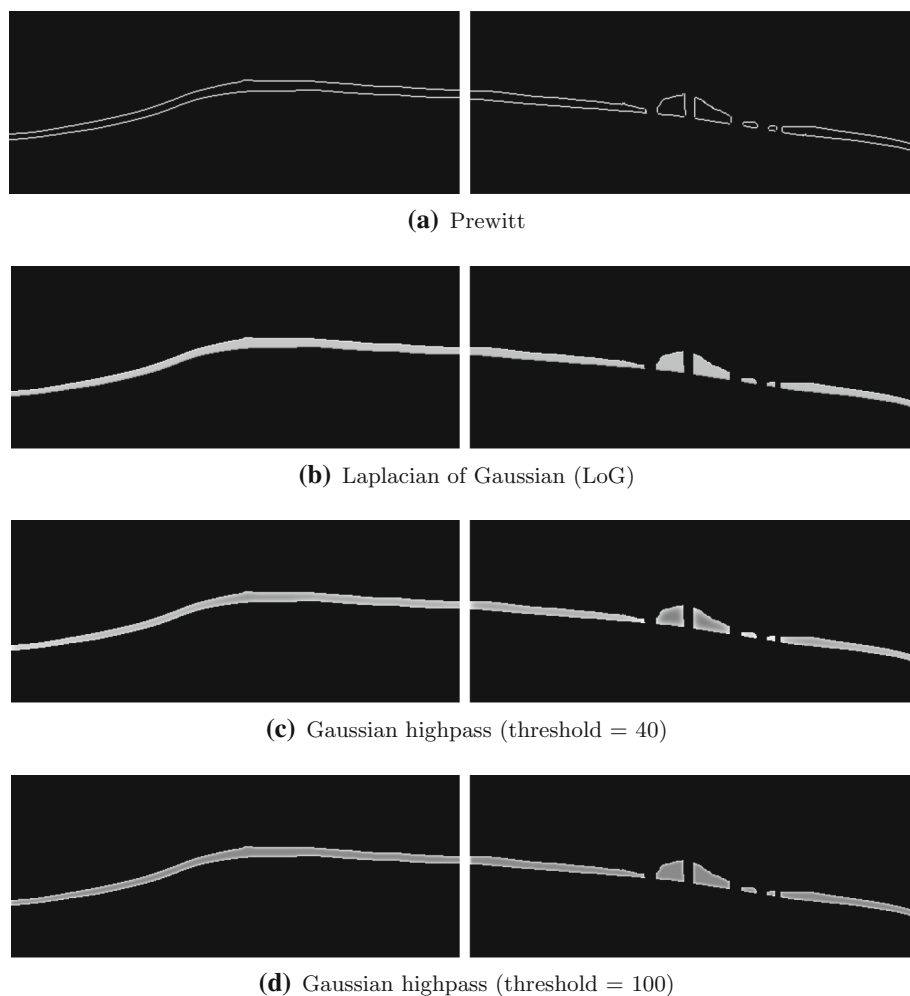
Bold values correspond to the highest AUC value, which serves us as a measure of accuracy

resentation. They are derived from convolutions with two edge filters, T_1 (Prewitt) and T_2 (Laplacian of Gaussian), and from two Gaussian high-pass filters, yielding T_3 and T_4 . Note that these feature transformations keep the same size, i.e., $T_j : \mathbb{R}^r \rightarrow \mathbb{R}^r$ for $j = 1, \dots, d$. See Fig. 6 for example images.

We can derive several augmented target loss functions L_p by choosing different $p \in \mathcal{G}_{k,d}$. In this experiment, we use the following projections:

- $p = I_4$,
- $\{p_l\}_{l=1}^{15}$, all projections from a t-design of strength 2 $\subset \mathcal{G}_{2,4}$ (see [10]),
- $p_{\text{PCA}} \in \mathcal{G}_{2,4}$, projection determined by PCA on the training data,
- $p_{\lambda_{2,4}}$, random projection chosen according to $\lambda_{2,4}$ in each mini-batch.

Fig. 6 Features on output and targets that enhance edges in different ways. It is not obvious which transformations are of most importance; weighting by projections can overcome this issue



6.3 Results

Since the detection problem is highly unbalanced, we use precision/recall curves [16] for evaluating the overall performance of each loss function model. The area under the curve (AUC) was used as a numerical indicator of the success rate [41]. The higher the AUC, the better the classification.

The results of the different loss functions on the independent test set are stated in Table 2. Due to the imbalance within the data, the photoreceptor region is identified well, but disruptions are not identified reliably when using the least squared loss function (5.1). For $\alpha = 0.1$, all proposed augmented target loss functions L_p clearly increase the success rate of the disruption quantification. Note that all projections are independent from the actual data set, except PCA that was computed beforehand on the training data.

The features itself (i.e., $p = I_4$) improve the quantification, and weighting them by projections increases the results even more: using the fixed projection p_{12} from the t-design sequence $\{p_l\}_{l=1}^{15}$ on the output/target features yields the highest accuracy for photoreceptors and disruptions. This

corresponds to the results of the previous sections, stating that depending on the particular data there are projections in the overall space acting beneficially. Since this projection generally cannot be found beforehand, using random projections in each loss function's evaluation step is easier, possible in practice, and independent from the data. The computation is efficient and randomization can have regularization effects that yield more robust results, cf. [34]. In the following, we will view a second classification problem based on spectrograms, where augmented target loss functions with random projections can improve the accuracy.

7 Application to Musical Data

Here, the learning task is a prototypical problem in Music Information Retrieval, namely multi-class classification of musical instruments. In analogy to the MNIST problem in image recognition, this classification problem is commonly used as a basis of comparison for innovative methods, since the ground truth is unambiguous and sufficiently many anno-

tated data are available. The input to the neural network is spectrograms of audio signals, which is the standard choice in audio machine learning. Spectrograms are calculated from the time signal using a short-time Fourier transform and taking the absolute value squared of the resulting spectra, thus yielding a vector for each time step and a two-dimensional array, like an image, cf. [18].

Reproducible code and more detailed information of our computational experiments can be found in the online repository [25].

7.1 Data and Objective

The publicly available GoodSounds data set [42] contains recordings of single notes and scales played by several single instruments. To gain equally balanced input classes, we restrict the classification problem to six instruments: clarinet, flute, trumpet, violin, alto saxophone and cello. Note that the recordings are monophonic, so that each recording yields one spectrogram that we aim to correctly assign to one of the six instruments.

After removing the silence [3,38], segments from the raw audio files are transformed into log-mel spectrograms [36], so that we obtain images of time–frequency representations with size 100×100 . One example spectrogram for each class of instruments is depicted in Fig. 7.

7.2 Convolutional Neural Network Learning

We implemented a fully convolutional neural network $f_\theta : \mathbb{R}^r \rightarrow [0, 1]^s$, cf. [33], where $r = 100 \times 100$ and $s = 6$, in Python 3.6 using Keras 2.2.4 framework [14] and trained it on the Nvidia GTX 1080 Ti GPU. The data are split into 1,40,722 training, 36,000 validation and 36,000 independent test samples. We heuristically choose $d = 16$ output features arising directly from the particular output class. The transformations $T_1, \dots, T_{16}$, with $T_j : \mathbb{R}^6 \rightarrow \mathbb{R}$ for $j = 1, \dots, 16$, are then given by the inner product of the output/target and the feature vectors. Among others, the features are chosen from the enhanced scheme of taxonomy [53] and from the table of frequencies, harmonics and under tones [57]. We use the proposed augmented target loss function L_p (5.2), where L_1 corresponds to the categorical cross-entropy

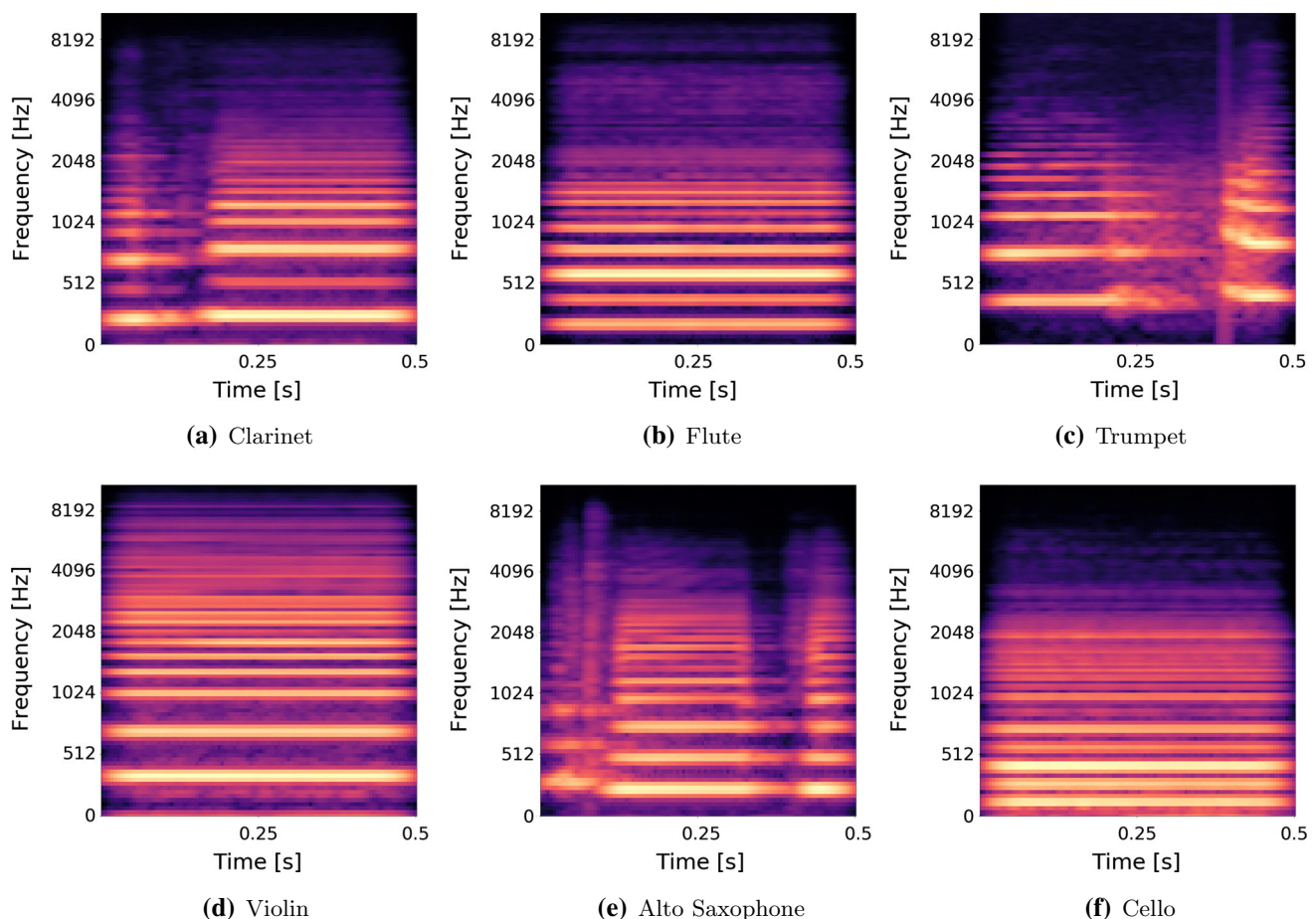


Fig. 7 Log-mel spectrograms of the six different instruments. Intensities range from 0 (black) to 1 (yellow) (Color figure online)

Table 3 Classification results with different parameter choices

α	β	p	Training	Test data
0	0	–	0.5541	0.5716
0.01	0	I_{16}	0.5650	0.5683
0.01	0	$p_{\lambda_{6,16}}$	0.7722	0.7657
0	0.05	–	0.9771	0.9729
0.01	0.05	I_{16}	0.9849	0.9802
0.01	0.05	$p_{\lambda_{6,16}}$	0.9857	0.9833

The standard inbuilt Tikhonov regularization (ℓ_2 -norm of θ) is weighted by β . For $\alpha > 0$, the feature transformations $\{T_j\}_{j=1}^{16}$ are used in the loss function, either directly or weighted by a random projection $p_{\lambda_{6,16}}$. The accuracy of the model is measured by the number of correctly classified samples divided by the number of all samples

Bold values correspond to the highest classification accuracy

loss [55] and L_2 to the mean squared error as in (5.3). We consider here two choices of p : the identity I_{16} and random projectors $p \sim \lambda_{6,16}$ in $\mathcal{G}_{6,16}$.

The deep learning model is sensitive to various hyper-parameters, including α and p , in addition to conventional parameters, such as the number of convolutional kernels, learning rate and the parameter β for Tikhonov regularization. To find the best choices in a fair trial, we utilize a random hyper-parameter search approach, where we train 60 models and select the three best ones for a more precise search over different α in the augmented target loss function and β for Tikhonov regularization. This results in 212 models that are evaluated on the training and validation set. Finally, we select the best model based on the accuracy of the validation set and evaluate it on the independent test set. For comparison, we also evaluate this model with no Tikhonov regularization, i.e., $\beta = 0$; see Table 3.

7.3 Results

Table 3 shows that no regularization and no features provide the poorest results. It seems that adding features with random projections has a regularizing effect and improves the results significantly. As expected, it is important to include Tikhonov regularization on θ . Further enhancement happens by adding features via the modified augmented target loss function with or without additional weighting from projections. All results are very stable and are generalizing very well from training to the independent test set; see [25] for further details.

Acknowledgements Open access funding provided by University of Vienna. This work was partially funded by the Vienna Science and Technology Fund (WWTF) through project VRG12-009, by WWTF AugUniWien/FA746A0249, by International Mobility of Researchers (CZ.02.2.69/0.0/0.0/16 027/0008371) and by project LO1401. For the research, infrastructure of the SIX Center was used.

Open Access This article is distributed under the terms of the Creative Commons Attribution 4.0 International License (<http://creativecommons.org/licenses/by/4.0/>), which permits unrestricted use, distribution, and reproduction in any medium, provided you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license, and indicate if changes were made.

Appendix A: Proof of Theorem 2.5

A.1 Proof of (2.14) in Theorem 2.5

For $\{y_i\}_{i=1}^M \subset \mathbb{R}^d$ and $p \in \mathcal{G}_{k,d}$, we define

$$f(p, \{y_i\}_{i=1}^M) := \frac{1}{M} \sum_{i=1}^M \frac{d}{k} \|p(y_i)\|^2. \quad (\text{A.1})$$

Given two sets, $\{y_i\}_{i=1}^{M_1}, \{z_j\}_{j=1}^{M_2} \subset \mathbb{R}^d$, suppose that $P \in \mathcal{G}_{k,d}$ is a random matrix, distributed according to a cubature measure of strength at least 2. The covariance is given by

$$\begin{aligned} & \text{Cov}(f(P, \{y_i\}_{i=1}^{M_1}), f(P, \{z_j\}_{j=1}^{M_2})) \\ &= \mathbb{E}[(f(P, \{y_i\}) - \mathbb{E}[f(P, \{y_i\})])(f(P, \{z_i\}) \\ & \quad - \mathbb{E}[f(P, \{z_i\})]) \end{aligned}$$

Using the identity, cf. [2],

$$\frac{d}{k} \mathbb{E}[\|Py\|^2] = \|y\|^2 \quad (\text{A.2})$$

directly yields

$$\begin{aligned} & \text{Cov}(f(P, \{y_i\}_{i=1}^{M_1}), f(P, \{z_j\}_{j=1}^{M_2})) \\ &= \mathbb{E} \left[\left(\frac{1}{M_1} \sum_{i=1}^{M_1} \frac{d}{k} \|P(y_i)\|^2 - \frac{1}{M_1} \sum_{i=1}^{M_1} \|y_i\|^2 \right) \right. \\ & \quad \left. \left(\frac{1}{M_2} \sum_{i=1}^{M_2} \frac{d}{k} \|P(z_i)\|^2 - \frac{1}{M_2} \sum_{i=1}^{M_2} \|z_i\|^2 \right) \right]. \end{aligned}$$

Following [8, Theorem 2.4, Sect. 3.1], we use that

$$\mathbb{E}[\|Py\|^2 \|Pz\|^2] = \frac{1}{q} (\alpha_1 \|y\|^2 \|z\|^2 + \alpha_2 \langle y, z \rangle^2), \quad y, z \in \mathbb{R}^d, \quad (\text{A.3})$$

holds, where $q = (d-1)d(d+2)$, $\alpha_1 = (d+1)k^2 - 2k$ and $\alpha_2 = 2k(d-k)$. This leads to the explicit formula of the population covariance

$$\begin{aligned} & \text{Cov} \left(f \left(P, \{y_i\}_{i=1}^{M_1} \right), f \left(P, \{z_j\}_{j=1}^{M_2} \right) \right) \\ &= \frac{a_{k,d}}{M_1 M_2} \sum_{i=1}^{M_1} \sum_{j=1}^{M_2} \langle y_i, z_j \rangle^2 \end{aligned}$$

$$- \frac{a_{k,d}}{d} \left(\frac{1}{M_1} \sum_{i=1}^{M_1} \|y_i\|^2 \right) \left(\frac{1}{M_2} \sum_{j=1}^{M_2} \|z_j\|^2 \right), \quad (\text{A.4})$$

with $a_{k,d} = \frac{2d(d-k)}{k(d-1)(d+2)}$.

For $y := \{y_i\}_{i=1}^M \subset \mathbb{R}^d \setminus \{0\}$, we set $\hat{y}_i := \frac{y_i}{\|y_i\|}$, for $i = 1, \dots, M$. The identity (A.4) enables us to compute the population correlation

$$\text{Corr}(f(P, y), f(P, \hat{y})) = \frac{\text{Cov}(f(P, y), f(P, \hat{y}))}{\sqrt{\text{Var}(f(P, y))} \sqrt{\text{Var}(f(P, \hat{y}))}} \quad (\text{A.5})$$

by the explicit formulas

$$\begin{aligned} \text{Cov}[f(P, y), f(P, \hat{y})] &= \frac{a_{k,d}}{M^2} \sum_{i,j=1}^M \langle y_i, \hat{y}_j \rangle^2 \\ &\quad - \frac{a_{k,d}}{d} \cdot \frac{1}{M} \sum_{i=1}^M \|y_i\|^2 \\ \text{Cov}[f(P, y), f(P, y)] &= \text{Var}[f(P, y)] = \frac{a_{k,d}}{M^2} \sum_{i,j=1}^M \langle y_i, y_j \rangle^2 \\ &\quad - \frac{a_{k,d}}{d} \left(\frac{1}{M} \sum_{i=1}^M \|y_i\|^2 \right)^2 \\ \text{Cov}[f(P, \hat{y}), f(P, \hat{y})] &= \text{Var}[f(P, \hat{y})] = \frac{a_{k,d}}{M^2} \sum_{i,j=1}^M \langle \hat{y}_i, \hat{y}_j \rangle^2 \\ &\quad - \frac{a_{k,d}}{d}. \end{aligned}$$

Since the variance is always nonnegative and $\frac{a_{k,d}}{d} > 0$, the denominator of $\text{Corr}(f(P, y), f(P, \hat{y}))$ in (A.5) satisfies

$$\begin{aligned} &\sqrt{\text{Var}(f(P, y))} \sqrt{\text{Var}(f(P, \hat{y}))} \\ &\leq \sqrt{\left(\frac{a_{k,d}}{M^2} \sum_{i,j=1}^M \langle y_i, y_j \rangle^2 \right) \left(\frac{a_{k,d}}{M^2} \sum_{i,j=1}^M \langle \hat{y}_i, \hat{y}_j \rangle^2 \right)} \\ &\leq \frac{a_{k,d}}{M^2} \sqrt{\left(\sum_{i,j=1}^M \langle y_i, y_j \rangle^2 \right) \left(\frac{1}{\min_i (\|y_i\|)^4} \sum_{i,j=1}^M \langle y_i, y_j \rangle^2 \right)} \\ &\leq \frac{1}{\min_i (\|y_i\|)^2} \frac{a_{k,d}}{M^2} \sum_{i,j=1}^M \langle y_i, y_j \rangle^2. \end{aligned}$$

The enumerator of $\text{Corr}(f(P, y), f(P, \hat{y}))$ in (A.5) is estimated by

$$\begin{aligned} \text{Cov}(f(P, y), f(P, \hat{y})) &\geq \frac{a_{k,d}}{\max_i (\|y_i\|)^2} \frac{1}{M^2} \sum_{i,j=1}^M \langle y_i, y_j \rangle^2 \\ &\quad - \frac{a_{k,d}}{d} \max_i (\|y_i\|)^2. \end{aligned}$$

For $d \geq M$, a short calculation yields $\text{Cov}(f(P, y), f(P, \hat{y})) \geq 0$, so that we obtain

$$\begin{aligned} \text{Corr}(f(P, y), f(P, \hat{y})) &\geq \frac{\min_i (\|y_i\|)^2}{\max_i (\|y_i\|)^2} \\ &\quad - \frac{\min_i (\|y_i\|)^2 \max_i (\|y_i\|)^2}{\frac{d}{M^2} \sum_{i,j=1}^M \langle y_i, y_j \rangle^2}. \end{aligned}$$

The lower bound $\sum_{i,j=1}^M \langle y_i, y_j \rangle^2 \geq M \min_i (\|y_i\|)^4$ yields

$$\begin{aligned} \text{Corr}(f(P, y), f(P, \hat{y})) &\geq \frac{\min_i (\|y_i\|)^2}{\max_i (\|y_i\|)^2} \\ &\quad - \frac{M}{d} \cdot \frac{\max_i (\|y_i\|)^2}{\min_i (\|y_i\|)^2}. \end{aligned}$$

Since the correlation is scaling invariant, the choice $y = \{x_i - x_j : 1 \leq i < j \leq m\}$ with $M = \frac{m(m-1)}{2}$ implies (2.14) in Theorem 2.5. Incorporating the correct scaling yields the following corollary:

Corollary A.1 For a given data set $x = \{x_i\}_{i=1}^m$ and for random $P \in \mathcal{G}_{k,d}$ the (co)variances of $\text{tvar}(Px)$ (2.4) and $\mathcal{M}(P, x)$ (2.8) are given by

$$\begin{aligned} \text{Cov}(\mathcal{M}(P, x), \text{tvar}(Px)) &= \frac{k}{2d} \left(\frac{a_{k,d}}{M^2} \sum_{i < j} \sum_{l < r} \left\langle x_i - x_j, \frac{x_l - x_r}{\|x_l - x_r\|} \right\rangle^2 \right. \\ &\quad \left. - \frac{a_{k,d}}{d} \cdot \frac{1}{M} \sum_{i < j} \|x_i - x_j\|^2 \right), \\ \text{Var}(\text{tvar}(Px)) &= \frac{k^2}{4d^2} \left(\frac{a_{k,d}}{M^2} \sum_{i < j} \sum_{l < r} \langle x_i - x_j, x_l - x_r \rangle^2 \right. \\ &\quad \left. - \frac{a_{k,d}}{d} \left(\frac{1}{M} \sum_{i < j} \|x_i - x_j\|^2 \right)^2 \right), \\ \text{Var}(\mathcal{M}(P, x)) &= \frac{a_{k,d}}{M^2} \sum_{i < j} \sum_{l < r} \left\langle \frac{x_i - x_j}{\|x_i - x_j\|}, \frac{x_l - x_r}{\|x_l - x_r\|} \right\rangle^2 \\ &\quad - \frac{a_{k,d}}{d}, \end{aligned}$$

where $M = \frac{m(m-1)}{2}$ and $a_{k,d} = \frac{2d(d-k)}{k(d-1)(d+2)}$.

A.2 Proof of the Second Part of Theorem 2.5

For fixed parameters $\mu > 0$, $\sigma^2 > 0$, that do not depend on d , let $Y_1 \in \mathbb{R}^d$ be a random vector, whose squared entries are independent, identically distributed with mean $\mathbb{E}Y_{1,l}^2 = \mu$ and variance $\text{Var}(Y_{1,l}^2) = \sigma^2$, for $l = 1, \dots, d$. We immediately observe

$$\mathbb{E}\left(\frac{\|Y_1\|^2}{\sqrt{d}}\right) = \sqrt{d}\mu, \quad \text{Var}\left(\frac{\|Y_1\|^2}{\sqrt{d}}\right) = \sigma^2.$$

For any $c > 0$, Chebyshev's inequality yields

$$\mathbb{P}\left(\left|\frac{\|Y_1\|^2}{\sqrt{d}} - \sqrt{d}\mu\right| \geq c\sigma\right) \leq \frac{1}{c^2}.$$

Suppose that $Y_2, \dots, Y_M$ are copies of Y_1 , not necessarily independent. Then, the union bound

$$\mathbb{P}\left(\left|\frac{\|Y_i\|^2}{\sqrt{d}} - \sqrt{d}\mu\right| \geq c\sigma, \text{ for some } i = 1, \dots, M\right) \leq \frac{M}{c^2}$$

implies that

$$\sqrt{d}\mu - c\sigma \leq \frac{\min_i (\|Y_i\|)^2}{\sqrt{d}} \leq \frac{\max_i (\|Y_i\|)^2}{\sqrt{d}} \leq \sqrt{d}\mu + c\sigma$$

holds with probability at least $1 - \frac{M}{c^2}$. Provided that $\sqrt{d}\mu \neq c\sigma$ and $0 < \sqrt{d}\mu - c\sigma$, we deduce

$$\frac{\sqrt{d}\mu - c\sigma}{\sqrt{d}\mu + c\sigma} \leq \frac{\min_i (\|Y_i\|)^2}{\max_i (\|Y_i\|)^2} \leq \frac{\sqrt{d}\mu + c\sigma}{\sqrt{d}\mu - c\sigma}.$$

We can choose $c = \frac{\mu}{\sigma} \sqrt{d}$, since $0 < c \leq \frac{\sqrt{d}\mu}{\sigma} \leq \frac{\sqrt{d}\mu}{\sigma}$. That directly yields

$$\frac{1 - \frac{1}{\sqrt{d}}}{1 + \frac{1}{\sqrt{d}}} \leq \frac{\min_i (\|Y_i\|)^2}{\max_i (\|Y_i\|)^2} \leq \frac{1 + \frac{1}{\sqrt{d}}}{1 - \frac{1}{\sqrt{d}}}$$

and holds with probability at least $1 - \frac{\mu^2 M}{\sigma^2 \sqrt{d}}$.

It follows directly that $\frac{\min_i (\|Y_i\|)^2}{\max_i (\|Y_i\|)^2}$ converges toward 1 in probability for $d \rightarrow \infty$,

The choice $\{Y_1, \dots, Y_M\} = \{X_i - X_j : 1 \leq i < j \leq m\}$ implies the second part of Theorem 2.5.

A.3 Calculations for Population Covariances

We notice that $\|p(x_i - x_j)\|^2 = \text{trace}(p x_i x_i^\top - p x_j x_j^\top)$ is a polynomial of degree 1 in p . Hence, $\text{tvar}(p x)$ in (2.4) is also a polynomial of degree 1 in p . If $\{p_l\}_{l=1}^n$ is a 1-design, then the sample mean of $\{\text{tvar}(p_1 x), \dots, \text{tvar}(p_n x)\}$ satisfies

$$\frac{1}{n} \sum_{l=1}^n \text{tvar}(p_l x) = \mathbb{E} \text{tvar}(P x),$$

which is the population mean of $\text{tvar}(P x)$, with $P \sim \lambda_{k,d}$. Similarly, the term $\|p(x_i - x_j)\|^4$ is a polynomial of degree 2 in p , so that $(\mathcal{M}(p, x))^2$ in (2.8) is a polynomial of degree 2 in p . If $\{p_l\}_{l=1}^n$ is a 2-design, then we derive

$$\begin{aligned} \sum_{l=1}^n (\mathcal{M}(p_l, x))^2 - \left(\sum_{j=1}^n \mathcal{M}(p_l, x) \right)^2 \\ = \mathbb{E}(\mathcal{M}(P, x))^2 - \mathbb{E} \left(\sum_{j=1}^n \mathcal{M}(P, x) \right)^2, \end{aligned}$$

with $P \sim \lambda_{k,d}$. In other words, the sample variance of $\{\mathcal{M}(p_1, x), \dots, \mathcal{M}(p_n, x)\}$ coincides with the population variance $\text{Var}(\mathcal{M}(P, x))$. Analogously, we deduce that the sample covariance of (3.5) coincides with the population covariance $\text{Cov}(\mathcal{M}(P, x), \text{tvar}(P x))$ with $P \sim \lambda_{k,d}$.

References

- Achlioptas, D.: Database-friendly random projections: Johnson–Lindenstrauss with binary coins. *J. Comput. Syst. Sci.* **66**(4), 671–687 (2003)
- Bachoc, C., Ehler, M.: Tight p -fusion frames. *Appl. Comput. Harmon. Anal.* **35**(1), 1–15 (2013)
- Bagwell, C.: SoX: Sound eXchange the Swiss army knife of sound processing. <https://launchpad.net/ubuntu/+source/sox/14.4.1-5>. Accessed 31 Oct 2018
- Ball, K.: An elementary introduction to modern convex geometry. *Flavors Geom.* **31**, 1–58 (1997)
- Baraniuk, R.G., Wakin, M.B.: Random projections of smooth manifolds. *Found. Comput. Math.* **9**, 941–944 (2006)
- Baraniuk, R., Davenport, M., DeVore, R., Wakin, M.: A simple proof of the restricted isometry property for random matrices. *Constr. Approx.* **28**(3), 253–263 (2008)
- Bingham, E., Mannila, H.: Random projection in dimensionality reduction: applications to image and text data. In: *Proceedings of the Seventh ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (New York, NY, USA), KDD '01, ACM, pp. 245–250 (2001)
- Bodman, B., Ehler, M., Gräf, M.: From low to high-dimensional moments without magic. *J. Theor. Probab.* **31**(4), 2167–2193 (2017)
- Breger, A., Ehler, M., Bogunovic, H., Waldstein, S.M., Philip, A., Schmidt-Erfurth, U., Gerendas, B.S.: Supervised learning and dimension reduction techniques for quantification of retinal fluid in optical coherence tomography images, *Eye*, Springer Nature (2017)
- Breger, A., Ehler, M., Gräf, M.: Quasi Monte Carlo Integration and Kernel-based Function Approximation on Grassmannians, *Frames and Other Bases in Abstract and Function Spaces*, Applied and Numerical Harmonic Analysis Series (ANHA). Springer, Birkhäuser (2017)
- Breger, A., Ehler, M., Gräf, M.: Points on manifolds with asymptotically optimal covering radius. *J. Complex.* **48**, 1–14 (2018)
- Candès, E.J., Tao, T.: Decoding by linear programming. *IEEE Trans. Inf. Theory* **51**(12), 4203–4215 (2005)
- Caruana, R.: Multitask learning. *Mach. Learn.* **28**(1), 41–75 (1997)
- Chollet, F., et al.: Keras (2015) <https://keras.io>
- Dasgupta, S., Gupta, A.: An elementary proof of a theorem of Johnson and Lindenstrauss. *Random Struct. Algorithms* **22**(1), 60–65 (2003)
- Davis, J., Goadrich, M.: The relationship between precision-recall and roc curves. In: *Proceedings of the 23rd International Conference on Machine Learning* (New York, NY, USA), ICML '06, ACM, pp. 233–240 (2006)
- de la Harpe, P., Pache, C.: Cubature formulas, geometrical designs, reproducing kernels, and Markov operators, *Infinite groups: geometric, combinatorial and dynamical aspects* (Basel), vol. 248, Birkhäuser, pp. 219–267 (2005)
- Dörfler, M., Bammer, R., Grill, T.: Inside the spectrogram: convolutional neural networks in audio processing. In: *IEEE International Conference on Sampling Theory and Applications* (SampTA), pp. 152–155 (2017)

19. Dörfler, M., Grill, T., Bammer, R., Flexer, A.: Basic filters for convolutional neural networks applied to music: training or design. *Neural Comput. Appl.* 1–14 (2018)
20. Dosovitskiy, A., Brox, T.: Generating images with perceptual similarity metrics based on deep networks. In: *Proceedings of the 30th International Conference on Neural Information Processing Systems (USA), NIPS'16*, Curran Associates Inc., pp. 658–666 (2016)
21. Etayo, U., Marzo, J., Ortega-Cerdà, J.: Asymptotically optimal designs on compact algebraic manifolds. *J. Monatsh. Math.* **186**(2), 235–248 (2018)
22. Ganaye, P.-A., Sdika, M., Benoit-Cattin, H.: Semi-supervised learning for segmentation under semantic constraint. In: *21st International Conference, Granada, Spain, Sept 16–20, 2018, Proceedings, Part III*, pp. 595–602 (2018)
23. Gerendas, B.S., Hu, X., Kaider, A., Montuoro, A., Sadeghipour, A., Waldstein, S.M., Schmidt-Erfurth, U.: Oct biomarkers predictive for visual acuity in patients with diabetic macular edema. *Investig. Ophthalmol. Vis. Sci.* **58**(8), 2026–2026 (2017)
24. Golub, G.H., Van Loan, C.F.: *Matrix Computations*, Johns Hopkins Studies in the Mathematical Sciences. The Johns Hopkins University Press (1996)
25. Harar, P.: Orthovar (2018) <https://gitlab.com/hararticles/orthovar>
26. Heckel, R., Tschannen, M., Böleskei, H.: Dimensionality-reduced subspace clustering. *Inf. Inference: J. IMA* **6**, 246–283 (2017)
27. Hedge, C., Sankaranarayanan, A.C., Yin, W., Baraniuk, R.G.: Numax: a convex approach for learning near-isometric linear embeddings. *IEEE Trans. Signal Process.* **83**, 6109–6121 (2015)
28. Johnson, J., Alahi, A., Fei-Fei, L.: Perceptual losses for real-time style transfer and super-resolution. In: *ECCV* (2016)
29. Karr, J.R., Martin, T.E.: Random numbers and principal components: further searches for the unicorn, Tech. report, United States Forest Service General Technical Report (1981)
30. Krahmer, F., Ward, R.: New and improved Johnson Lindenstrauss embeddings via the restricted isometry property. *SIAM J. Math. Anal.* **43**(3), 1269–1281 (2011)
31. Ledig, C., Theis, L., Huszar, F., Caballero, J., Aitken, A.P., Tejani, A., Totz, J., Wang, Z., Shi W.: Photo-realistic single image super-resolution using a generative adversarial network. In: *2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 105–114 (2017)
32. Linial, N., London, E., Rabinovich, Y.: The geometry of graphs and some of its algorithmic applications. *Combinatorica* **15**(2), 215–245 (1995)
33. Long, J., Shelhamer, E., Darrell, T.: Fully convolutional networks for semantic segmentation. In: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 3431–3440 (2015)
34. Mahoney, M.W.: Randomized algorithms for matrices and data. *Found. Trends@ Mach. Learn.* **3**(2), 123–224 (2011)
35. Matousek, J.: On variants of the Johnson–Lindenstrauss lemma. *Random Struct. Algorithms* **33**(2), 142–156 (2008)
36. McFee, B., et al.: Librosa: 0.6.2 (2018) <https://doi.org/10.5281/zenodo.1342708>
37. Moore, B.: Principal component analysis in linear systems: controllability, observability, and model reduction. *IEEE Trans. Autom. Control* **26**, 17–32 (1981)
38. Navarrete, J.: The sox of silence (2009) <https://digitalcardboard.com/blog/2009/08/25/the-sox-of-silence>
39. Neumayer, S., Nimmer, M., Setzer, S., Steidl, G.: On the robust PCA and Weiszfeld's algorithm. *Appl. Math. Optim.* 1–32 (2019)
40. Oktay, O., Ferrante, E., Kamnitsas, K., Heinrich, M., Bai, W., Caballero, J., Guerrero, R., Cook, S.A., de Marvao, A., Dawes, T., O'Regan, D., Kainz, B., Glocker, B., Rueckert, D.: Anatomically constrained neural networks (ACNN): application to cardiac image enhancement and segmentation. *IEEE Trans. Med. Imaging* **37**, 384–395 (2018)
41. Pabst, G.: Parameters for Compartment-Free Pharmacokinetics-Standardisation of Study Design, Data Analysis and Reporting, ch. 5. Area Under the Concentration-Time Curve, pp. 65–80, Shaker Verlag (1999)
42. Picas, O.R., Rodriguez, H.P., Dabiri, D., Tokuda, H., Hariya, W., Oishi, K., Serra, X.: A real-time system for measuring sound goodness in instrumental sounds. In: *Audio Engineering Society Convention 138*, Audio Engineering Society (2015)
43. Ronneberger, O., Fischer, P., Brox, T.: U-net: convolutional networks for biomedical image segmentation (2015) [arXiv:1505.04597](https://arxiv.org/abs/1505.04597)
44. Seymour, P., Zaslavsky, T.: Averaging sets: a generalization of mean values and spherical designs. *Adv. Math.* **52**, 213–240 (1984)
45. Simonyan, K., Zisserman, A.: Very deep convolutional networks for large-scale image recognition (2014) [arXiv:1409.1556](https://arxiv.org/abs/1409.1556)
46. Stauffer, D.F., Garton, E.O., Steinhorst, R.K.: Ecology: a comparison of principal components from real and random data. *Ecology* **66**(6), 1693–1698 (1985)
47. Stewart, R., Ermon, S.: Label-free supervision of neural networks with physics and domain knowledge. In: *AAAI* (2017)
48. Thanei, G.-A., Heinze, C., Meinshausen, N.: Random Projections for Large-scale Regression, pp. 51–68. Springer, Cham (2017)
49. Udell, M.: Generalized low rank models, Ph.D. thesis, Stanford University (2015)
50. Varnousfaderani, E.S., Wu, J., Vogl, W.-D., Philip, A.-M., Montuoro, A., Leitner, R., Simader, C., Waldstein, S.M., Gerendas, B.S., Schmidt-Erfurth, U.: A novel benchmark model for intelligent annotation of spectral-domain optical coherence tomography scans using the example of cyst annotation. *Comput. Methods Programs Biomed.* **130**, 93–105 (2016)
51. Veraart, Jelle, Novikov, Dmitry S., Christiaens, Daan, Ades-aron, Benjamin, Sijbers, Jan, Fieremans, Els: Denoising of diffusion MRI using random matrix theory. *NeuroImage* **142**, 394–406 (2016)
52. Vershynin, R.: Introduction to the non-asymptotic analysis of random matrices. In: Eldar, Y., Kutyniok, G. (eds.) *Compressed Sensing, Theory and Applications*, pp. 210–268. Cambridge University Press, Cambridge (2012)
53. von Hornbostel, E.M., Sachs, C.: Classification of musical instruments: translated from the original german by anthony baines and klaus p. wachsmann. *Galpin Soc. J.* 3–29 (1961)
54. Yosinski, J., Clune, J., Bengio, Y., Lipson, H.: How transferable are features in deep neural networks? In: *Proceedings of the 27th International Conference on Neural Information Processing Systems—Volume 2* (Cambridge, MA, USA), NIPS'14, MIT Press, pp. 3320–3328 (2014)
55. Zhang, Z., Sabuncu, M.: Generalized cross entropy loss for training deep neural networks with noisy labels. In: *Advances in Neural Information Processing Systems (NIPS)* **31** (2018)
56. Zhang, L., Lukac, R., Wu, X., Zhang, D.: Pca-based spatially adaptive denoising of cfa images for single-sensor digital cameras. *IEEE Trans. Image Process.* **18**(4), 797–812 (2009)
57. ZyTrax Inc.: Frequency ranges (2018) <http://www.zytrax.com/tech/audio/audio.html>



ical School, Boston (MA).

A. Breger received her MSc degree in Mathematics from the University of Vienna in 2015, followed by a research position in collaboration with the Medical University of Vienna. She focuses on interdisciplinary research with her PhD studies on 'Dimension reduction with orthogonal projections and applications on medical images with learning tasks'. In 2019 she received a grant to work at the Laboratory of Mathematics in Imaging (LMI), Brigham and Women's Hospital, Harvard Med-



ted to music applications.

M. Dörfler obtained her PhD in Mathematics from the University of Vienna and is a research assistant at the Mathematics Department of the University of Vienna. She is heading her own research group and is interested in the interplay between aspects of harmonic analysis and the design of learning algorithms based on convolutional neural networks. She has worked on audio signal processing and, being a musician herself, has been particularly interested in developing representations adap-



Yatiris from Pladema Institute (UNICEN). His primary research interests include machine learning and computer vision applied to computer-assisted ophthalmology and medicine in general.

J. I. Orlando obtained his Ph.D. in computational and industrial mathematics from Universidad Nacional del Centro de la Provincia de Buenos Aires (UNICEN), Tandil, Argentina. In 2018 and 2019, he has worked as a Postdoctoral Researcher at the Department of Ophthalmology of the Medical University of Vienna (Austria). He is currently a permanent researcher from the National Scientific and Technical Research Council (CONICET, Argentina), working at the research group



S. Klimescha is a team member of the OPTIMA lab and a resident at the Department of Ophthalmology, Medical University of Vienna. She obtained her MD degree at the Medical University of Vienna, Austria in July 2015, joined the OPTIMA team in January 2016 and is enrolled in the PhD Program for Medical Imaging.



P. Harar obtained a M.Sc. in System Engineering and Informatics from Brno University of Technology in 2015 before pursuing a Ph.D. at the Brain Diseases Analysis Laboratory on "Audio Classification with Deep Learning on Limited Data Sets". He has been working at the Numerical Harmonic Analysis Group (University of Vienna) since 2018. His research focuses on medical audio analysis using deep learning methods and their development.



C. Grechenig obtained his MD degree at the Medical University of Graz, Austria in January 2017, and joined the Christian Doppler Laboratory for Ophthalmic Image Analysis (OPTIMA) in March 2018 and is enrolled in the PhD Program for Medical Imaging at the Medical University of Vienna. As a clinician he currently works as a retina resident at the Department of Ophthalmology and Optometry, Medical University of Vienna.



B. S. Gerendas is the Managing Director of the Vienna Reading Center at the Department of Ophthalmology, Medical University of Vienna since 2017. She completed Medical School in 2009 at the University of Heidelberg, Germany, where she in parallel obtained a Master of Science degree in Healthcare Management. Besides her PhD in Medical Physics, she is a board certified ophthalmologist/retina specialist and finished her habilitation (Assoc. Prof.) in 2018. Her primary

research interests are in the field of retinal and choroidal imaging and (automated) image analysis as well as translational ophthalmic research in retinal diseases.



U. Schmidt-Erfurth is Chair of the Department of Ophthalmology at the Medical University of Vienna, Austria. Her areas of expertise include surgical and medical retina and the development of innovative diagnostic techniques and treatment strategies for chorioretinal diseases. She is the Head of the Vienna Reading Center and the founder and leader of the OPTIMA project which introduces means of artificial intelligence into retinal imaging. She is the author of over 420 original articles. She has

received numerous grants/awards and is a member of the American Academy of Ophthalmology, EURETINA, the European and the Austrian Academy of Sciences.



M. Ehler After his PhD in mathematics in 2007, he received a 3 year PostDoctoral Fellowship at the National Institutes of Health in Bethesda/Maryland. His research efforts are guided by bridging gaps between theoretical and numerical mathematics and medical applications. Since 2013 he is the head of a Vienna Research Group for Young Investigators on computational harmonic analysis of high-dimensional biomedical data. At the end of 2017 he became Assoc. Professor at the

Faculty of Mathematics at the University of Vienna.

(P3)

Supervised Learning and Dimension Reduction Techniques for Quantification of Retinal Fluid in Optical Coherence Tomography Images

Journal: Eye, Springer Nature

Authors: A. Breger, M. Ehler, H. Bogunovic, S. M. Waldstein, A. Philip, U. Schmidt-Erfurth, and B. S. Gerendas

Published: 2017

Link: <https://www.nature.com/articles/eye201761>

DOI: 10.138

Contribution: Conceptual idea; Design and performance of the numerical experiments; Main writing.

Supervised learning and dimension reduction techniques for quantification of retinal fluid in optical coherence tomography images

A Breger¹, M Ehler¹, H Bogunovic²,
SM Waldstein², A-M Philip², U Schmidt-Erfurth²
and BS Gerendas²

Abstract

Purpose The purpose of the present study is to develop fast automated quantification of retinal fluid in optical coherence tomography (OCT) image sets.

Methods We developed an image analysis pipeline tailored towards OCT images that consists of five steps for binary retinal fluid segmentation. The method is based on feature extraction, pre-segmentation, dimension reduction procedures, and supervised learning tools.

Results Fluid identification using our pipeline was tested on two separate patient groups: one associated to neovascular age-related macular degeneration, the other showing diabetic macular edema. For training and evaluation purposes, retinal fluid was annotated manually in each cross-section by human expert graders of the Vienna Reading Center. Compared with the manual annotations, our pipeline yields good quantification, visually and in numbers.

Conclusions By demonstrating good automated retinal fluid quantification, our pipeline appears useful to expert graders within their current grading processes. Owing to dimension reduction, the actual learning part is fast and requires only few training samples. Hence, it is well-suited for integration into actual manufacturer's devices, further improving segmentation by its use in daily clinical life.

Eye advance online publication, 21 April 2017; doi:10.1038/eye.2017.61

Introduction

Spectral-domain optical coherence tomography (OCT) is nowadays the most frequently used

retinal imaging technique in ophthalmology. It provides good visualization of subretinal and intraretinal cystoid fluid. Especially distinguishing between fluid and no fluid or a fluid increase or decrease between two visits is important for the administration of anti-vascular endothelial growth factor therapy. However, automated quantification of retinal fluid in OCT images goes beyond evaluation algorithms provided by manufacturers. When considering that a regular OCT raster scan has between 128 and 1024 B-scans, this can be tedious and beyond possibilities of every day clinical routine (cf. Schmidt-Erfurth *et al.*¹).

Owing to increased computing power, state-of-the-art methods for several learning tasks related to image analysis in the applied sciences are deep neural networks (NNs).² Deep convolutional NNs, in particular, have already been applied to fluid quantification in OCT images (cf. Schlegl³). Deep learning models have several drawbacks, such as the need for huge amounts of training data, high training time, large memory usage, and the complicated design of a suitable network. A semi-supervised technique was developed by Zheng *et al.*⁴ (see also Chen *et al.*,⁵ Xu *et al.*,⁶ and Wilkins *et al.*,⁷ for general learning techniques, we refer to Grady,⁸ Czaja and Ehler,⁹ Schölkopf and Smola,¹⁰ and Roweis and Saul¹¹). Here we shall provide an efficient method with a fast learning part and low memory needs based on dimension reduction techniques. Moreover, it shall be easy to implement, requires only few training data, and yields good fluid quantification. To meet clinical needs, we present a fully automated image analysis pipeline tailored to fluid quantification in OCT raster image sets. The

¹Department of Mathematics, University of Vienna, Vienna, Austria

²Vienna Reading Center and Christian Doppler Laboratory for Ophthalmic Image Analysis, Department of Ophthalmology, Medical University of Vienna, Vienna, Austria

Correspondence: BS Gerendas, Vienna Reading Center and Christian Doppler Laboratory for Ophthalmic Image Analysis, Department of Ophthalmology, Medical University of Vienna, Waehringer Guertel 18-20, Vienna 1090, Austria
Tel: +4314040079310;
Fax: +4314040079320.
E-mail: bianca.gerendas@meduniwien.ac.at

Received: 26 February 2017
Accepted in revised form: 15 March 2017

pipeline is built by five steps that are guided by the principle of dimension reduction. Our resulting binary segmentations are compared with the manual annotations of the expert graders visually and in numbers, illustrating the strengths and weaknesses of our approach. The innovative aspect of our approach is the combination of local and global features in dimension reduction techniques together with learning methods. Our approach presents a model applicable for clinical routine and could be implemented into a OCT device for its every day use.

Methods

The method is designed for a collection of 3D OCT volumes. For training and evaluation purposes, retinal fluid was annotated manually in each cross-section by human expert graders of the Vienna Reading Center (VRC). Guaranteeing unbiased evaluation, the data were annotated before the development of the pipeline and hence independently analyzed. As manual annotation is very time consuming (~15 h per volume), there are few training data available, that is, binary annotations of retinal fluid in few cross-sections, which enable the application of supervised learning techniques.

Our image analysis pipeline is built up from five steps. To reduce computation costs and minimize required training data, we focus on the extraction of the necessary image information contents before entering the learning stages. More specifically, (1) we compute local features for each B-scan that enhance specific image characteristics, while inheriting effective denoising. For each OCT B-scan, we compute two local features. One results from Gabor filtering¹² and the other one as the solution of a variational minimization problem¹³ resulting in a piecewise smooth image. These two local features form a vector-valued image (*cf.* Local features). (2) We replace each B-scan with the piecewise constant solution of another variational energy minimization problem,¹⁴ whose input is the vector-valued image. This solution is considered as a pre-segmentation, in which darker piecewise constant areas appear to be associated with retinal fluid (*cf.* Dimension reduction I: pre-segmentation). We now aim to select those areas by automated learning an appropriate threshold for each cross-section, such that segments with pixel values underneath the threshold represent the fluid. To minimize computation costs and reduce the required training data in the learning process, (3) we compute four global parameters of each cross-section (*cf.* Dimension reduction II). (4) We now apply three standard learning tools (regression, decision trees, and NNs) to estimate the thresholds from these parameters based on few randomly selected manual annotations of expert graders serving as training samples (*cf.* Supervised learning of the

thresholds). We then evaluate the trained learning models on the global parameters of each B-scan of the remaining test samples. (5) Once the thresholds are estimated, we use them to threshold the piecewise constant pre-segmentation images obtaining the final identification of retinal fluid areas (*cf.* Thresholding). Steps 2 and 3 are significant dimension reduction processes enabling faster learning with fewer training data. We run through these five steps for each 2D cross-sectional image independently, yielding the following analysis pipeline:

- Local features (pixel-level):
Compute two features for each B-scan, yielding vector-valued images.
- Dimension reduction I (pre-segmentation):
Use these vector-valued images to replace each B-scan with the piecewise constant solution of an energy minimization problem.
- Dimension reduction II:
Represent each B-scan by few global features (image level).
- Learning thresholds:
From these global features estimate the thresholds via standard learning tools.
- Thresholding:
Each piecewise constant B-scan is thresholded to obtain the final classification of the fluid spaces.

Local features

We extract two local features of each B-scan (see the first two images in Figure 1). One feature image is obtained from Gabor filtering (see Supplementary Equation 1). Gabor filters are well-suited for texture analysis;¹² see also Soares *et al*¹⁵ for applications in retinal image analysis. The second feature is obtained from the Rudin–Osher–Fatemi variational energy minimization¹³ (see Supplementary Equations 2 and 3). Its solution is a piecewise smooth image, that is, edges are preserved and regions between interfaces are smoothed.

Dimension reduction I: pre-segmentation

We now use the vector-valued feature images to replace each B-scan with the piecewise constant solution of another variational energy minimization problem, (see Supplementary Equation 4 and *cf.* Mumford and Shah,¹⁶ and Potts.¹⁷ We minimize a discrete domain version (see Supplementary Equation 5) by using the Pottslab matlab code from the website <http://www.pottslab.de> (*cf.* Storath and Weinmann¹⁴), where an alternating direction method of multipliers is used.

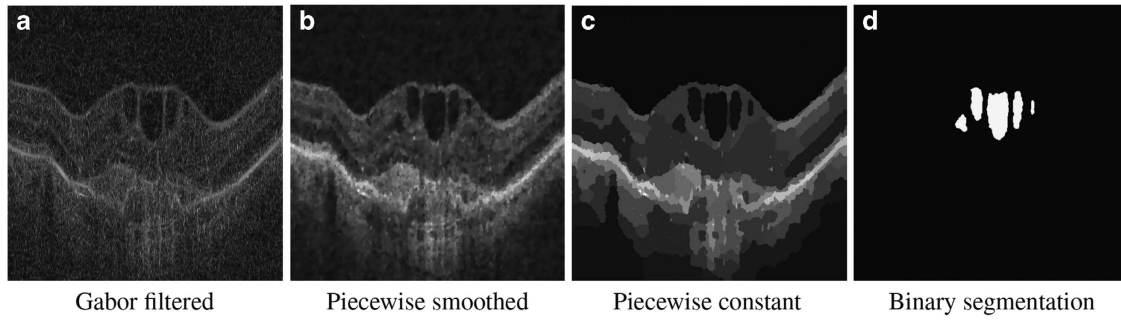


Figure 1 Example image showing results after Step 1 (a and b), 2 (c), and 5 (d) of the pipeline (described in Local features, Dimension reduction I: presegmentation, and Thresholding).

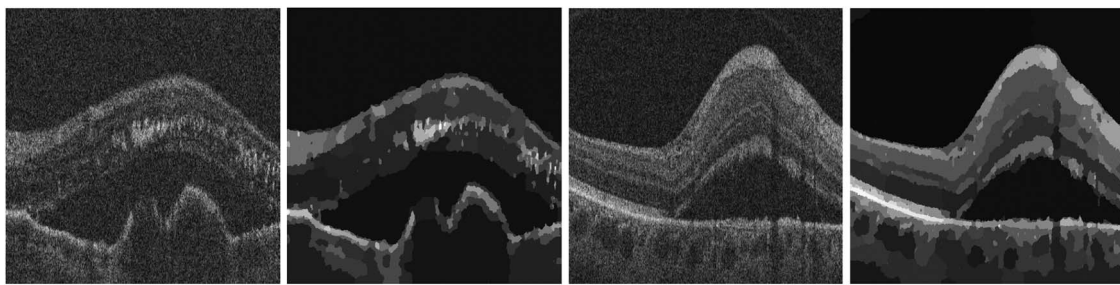


Figure 2 The first and third images show original cross-sections, and the second and fourth the corresponding piece-wise constant images, which are the basis for the supervised learning (see Dimension reduction I: pre-segmentation).

It is noteworthy that the minimizer is again a vector-valued image. We use the second channel of the piecewise constant image, which corresponds to the Rudin–Osher–Fatemi model, see Figure 2. Indeed it yields pre-segmentations in which retinal fluid appears as dark regions and boundaries seem preserved well. Thus, it seems reasonable to threshold the piecewise constant image to divide the piecewise constant images in two regions: white (presence of retinal fluid) and black (absence). As annotations from human graders are available, supervised learning techniques can be used to estimate the magnitude of appropriate thresholds.

Dimension reduction II

In order to learn the thresholds, we need to specify the parameters of the piecewise constant image that actually characterize the thresholds magnitude. We observed that the threshold is mainly influenced by few global image parameters (mean, variance, infimum, and median). Therefore, we associate just these four global features of the region of interest (ROI) with each piecewise constant image that are used as input for subsequent learning.

Supervised learning of the thresholds

We split our cross-sections into training and test samples. As training data, we use few randomly chosen 2D slices

and test it on all remaining cross-sections. First, we manually determine thresholds $t_i \in [0, 1]$ as target values, based on the optimal fit with the graders' annotations.

To estimate thresholds from four global features (see Dimension reduction II), we make use of three learning techniques: regression, random forest, and three-layer NNs (cf. Bishop,¹⁸ and Witten and Frank¹⁹ for more details). The target thresholds $\{t_i\}$ are used to train these learning models. We obtain the remaining thresholds for the test data set, which is disjoint from the training data, by evaluating the trained learning model on the global features of the test samples. The machine learning and NN toolbox of MatlabR2015a is used for the learning computations.

Thresholding

The last step of our pipeline is to threshold all the computed pre-segmentations (see Dimension reduction I: pre-segmentation) of the test set by the evaluated thresholds. Those segments with values less than the threshold are classified as retinal fluid (see right image in Figure 1).

Test data

Our proposed analysis pipeline is tested on imaging data provided by the VRC. The VRC has analyzed millions of images, mainly from OCT, recorded in multi-center clinical trials with many hundred different study sites all

Table 1 Evaluation of classification and volume quantification for the different learning methods and the random training data sets

Classification (image-wise)	ADM data				DME data			
	Regression	RF (12 trees)	NN (20 HU)	Training	Regression	RF (7 trees)	NN (5 HU)	Training
Correctly classified (mean)	72%	70%	73%	88%	75%	72%	76%	87%
Sensitivity	0.69	0.66	0.72	0.83	0.77	0.74	0.77	0.95
Specificity	0.75	0.74	0.74	0.95	0.63	0.67	0.70	0.65
Quantification (pixel-wise)								
Dice coefficient	0.54	0.50	0.53	0.63	0.41	0.47	0.50	0.59
Difference vol in mm ³ (mean, variance)	(0.22, 0.15)	(0.24, 0.10)	(0.24, 0.12)		(0.47, 0.18)	(0.25, 0.07)	(0.31, 0.03)	

Abbreviations: ADM, age-related macular degeneration; DME, diabetic macular edema; HU, hidden units; NN, neural network; RF, random forest.

over the world in a standardized manner over the past decade. We shall work here with two separate OCT data sets, the first corresponds to 38 patients suffering from neovascular age-related macular degeneration (AMD), where OCT images are taken at the very start of the treatment initiation (ie, all are baseline images). The second data set relates to 16 patients suffering from diabetic macular edema (DME) with different severity of edema selected randomly from a trial database within the VRC. All OCT images were recorded by a Cirrus OCT camera (Carl Zeiss Meditec, Dublin, CA, USA). We will test the pipeline on both data sets independently.

Results

For each patient of the two data sets, the recorded retinal image data yields a three-dimensional ($1024 \times 512 \times 128$) matrix, that is, the volume data consists of 128 cross-sectional images of one patient's retina of one eye each of the size (1024×512). We cut out scaled images of the size (512×512) that include the ROI, that is, the region between the inner limiting membrane and the pigment epithelium cells, which is the retina (see Garvin *et al*²⁰ and Haeker *et al*²¹ for retinal layer segmentation). This halves the amount of pixels of each 2D image without losing any information. Retinal fluid was annotated in each of the $4864 = 38 * 128$ and $2048 = 16 * 128$ cross-sections manually by human expert graders of the VRC, provided as binary image sets. Using professional drawing tablets, all annotations were performed B-scan-wise. The time required to annotate a typical OCT volume was ~15 h. In total, five readers and three board-certified supervisors, who corrected errors if required, participated in the annotation procedure. Each reader underwent a standardized training before annotating the study data set (see Waldstein *et al*²² and Varnousfaderani *et al*²³ for more details).

The results of our fully automated quantification of subretinal and intraretinal cystoid fluid highly depend on

the quality of the piece-wise constant pre-segmentation.

Table 1 shows good binary segmentation outcomes for both data sets when we use the optimal threshold on the piece-wise constant training data and compare it with the graders' annotations. It is noteworthy that sensitivity, specificity, and dice coefficients from the training sets (see Table 1) serve as the baselines for evaluations on the test data.

Qualitative/visual evaluation

In Figure 3, we show original OCT images on the left, next to the same images with the retinal fluid annotated by human expert graders in the center. On the right, we provide the output of our image analysis pipeline, which is our resulting binary segmentation. In these selected cases, retinal fluid locations are relatively unambiguous and graders' annotations are consistent with the visual impressions. Our segmentations coincide well with the graders' annotations.

The human expert graders annotate each cross-section separately, which leads to inconsistencies in few subsequent cross-sections that appear almost identical.

In Figure 4a, we see that similarly looking cross-sections are annotated differently by human graders, whereas our algorithm identifies retinal fluid consistently in both cases. Which is clinically correct in this case remains unclear, as the fluid can be hidden by other pathologic structures and the human annotations always leave a little room for subjectiveness and interpretation, even when high standards for objectivity are set. In any case, a consistent result is important and this can only be achieved by the algorithm. Retinal vessels cast shadows in OCT, attenuating the signal underneath and yielding darker regions in the original images. It should be noticed that, without further substeps involved, such regions get marked incorrectly by our pipeline (*cf.* Figure 4b).

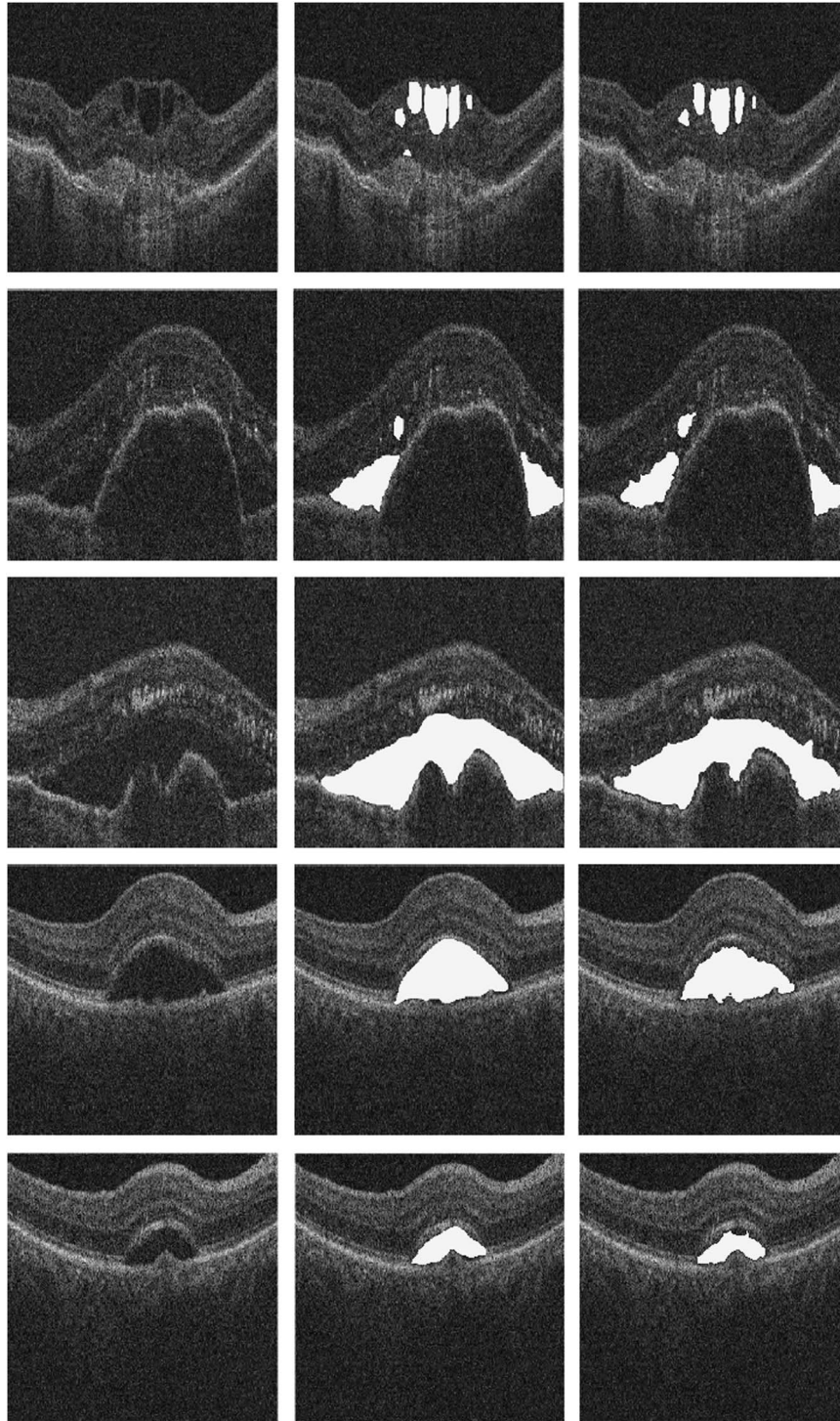


Figure 3 Final binary segmentation results are shown: (left) original image, (center) graders' annotations, and (right) pipeline output.

Quantitative evaluation

In Table 1, we compare our resulting binary segmentations with the manual annotations of the expert graders serving as the reference of quality.

Two types of quantitative evaluations are considered. We refer to classification when we aim to detect cross-sections that contain retinal fluid, that is, each 2D cross-section is globally classified into either presence or absence of retinal fluid. We refer to volume

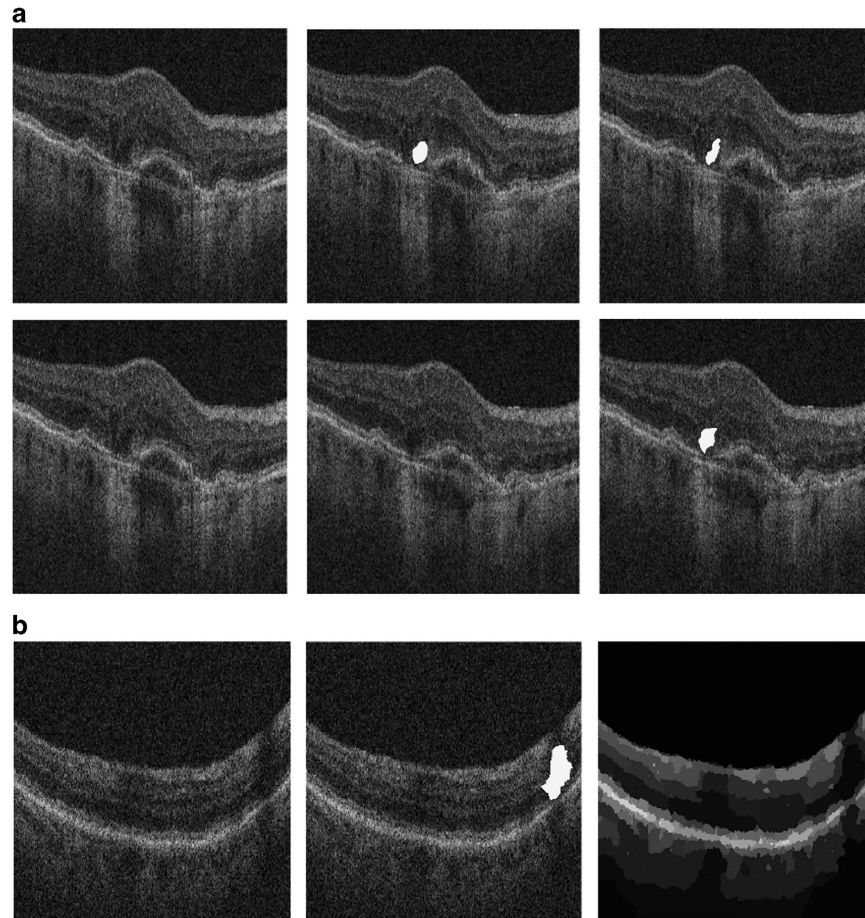


Figure 4 Some challenges of the processing: (a) similar looking images are annotated differently by human graders, whereas our algorithm identifies retinal fluid consistently in both cases. Row 1 and 2: (left) original image, (center) graders' annotations, and (right) pipeline output. (b) Shadows from vessels are marked wrong by our pipeline. Row 3: (left) original image, (center) graders' annotations, and (right) pre-segmentation.

quantification when comparing the pipelines output in the ROI pixel-wise with the graders' annotations in each cross-section. Classification and volume quantification are compared for the two data sets separately. Comparing the three different learning techniques sensitivity and specificity for classification are at the same level. The NNs are above the 70% bound for both data sets. For volume quantification, regression lacks behind on the DME data. NNs, on the other hand, yield good success rates with dice coefficients bigger than 0.5 for both data sets.

We also provide mean and variance difference between the manual annotations and our quantified retinal fluid volume in mm^3 , where camera specifications yield that one pixel translates to $\approx 10^{-6} \text{mm}^3$. The corresponding Bland–Altman plot in Figure 1 of the Supplementary Material shows reasonable agreement between annotations and computed segmentation for the volume quantification using the NNs.

Computational resources

The run-time in MatlabR2015a for one OCT volume (128 B-Scans) on a 2.2 GHz Intel Core i7 processor with 16 GB RAM is ~ 45 min. The vector-valued minimization (Supplementary Equation 5) is the most time-consuming part (43 min), which is independent from the actual learning taking less than a second.

Discussion

The development of our pipeline was guided by the idea of dimension reduction in order to simplify the actual learning part. Although the pre-segmentation is the first dimension reduction by significantly quantizing the intensity range of the image, the most extensive dimension reduction stems from learning the thresholds from four global features per cross-section. It seems that the threshold effectively only depends on global statistics

of the pre-segmentation, enabling cheap and fast learning. The second step covers the essential image information. Gabor filtering takes care of texture and the Rudin–Osher–Fatemi model enhances edges, while denoising between interfaces. The subsequent piecewise constant energy minimization yields a pre-segmentation and is the basis for final binary segmentation through thresholding. We have seen that this pre-segmentation is sufficiently rich, so that optimal thresholding on the training data yields regions that coincide well with the provided manual annotations. The pre-segmentation is time-consuming relative to the other pipeline steps. Faster pre-segmentation methods are available but the piecewise constant minimizer of (Supplementary Equation 4) yields the best results for our concerns at the moment.

There is a beneficial interaction between features, pre-segmentation, dimension reduction, and threshold learning, which yields convincing results. For each patient set, we used only 100 randomly selected B-scans as training samples for evaluating the pipeline on the remaining cross-sections yielding 4764 (AMD) and 1948 (DME) test samples. It is a big benefit that only few manually annotated training data are needed in our pipeline, as large amounts of manual annotations are often not available. Our image analysis pipeline for quantification of retinal fluid in OCT has the potential to significantly support human expert graders. By our automated classification into presence or absence of retinal fluid, graders can disregard a large subset of B-scans speeding up the manual grading process. Our results for the detection yield sufficient sensitivity and specificity rates, so that its integration into the grading process seems useful. Especially considering the growing number of B-scans and details in each B-scan due to higher resolution in all scan directions, automated analysis methods will be needed in ophthalmology. Here, it is of greatest importance to have semi-automated segmentation tools for the validation of proposed automated methods. Our method will speed up the annotation in a reading center setting by giving a fast first attempt as the basis for manual annotation correction, or within the frame of clinical routine, giving the ophthalmologist a quick first result for the interpretation of the scan. When compared with other commonly used computerized evaluation in biomedicine, such as automated interpretation of electrocardiograms, our pipeline appears definitely competitive (*cf.* NAMEstes²⁴ and Wetherell²⁵).

Concerning our automated quantification of retinal fluid, we saw in Results that visual comparison of our results with the graders' annotations appear convincing. For high-quality manual annotations, our segmented

regions coincide well. However, earlier studies have shown that inconsistencies between graders are normal to a certain percentage and we can only aim at the same inter-grader reproducibility such as the reported human inter-grader variability (*cf.* Varnousfaderani²³). In Figure 4a, cross-sections appear similar but graders' annotations differ. Medical experts confirm inconsistency and our automated approach delivers consistent segmentations in that case. Taking away subjectivity by using automated segmentation is of particular importance when it comes to tracking retinal fluid changes in patients over time. Such deviations from the annotations also contribute to the error rates in the quantitative evaluation (Table 1). Further falsely segmented areas are due to vessel shadows, which can easily be removed by human graders' when our pipeline is used as a supporting tool as indicated above. For a fully automated quantification, vessel shadows should be taken care of in additional processing steps (*cf.* Girard *et al*,²⁶ Wu *et al*,²⁷ and Ehler *et al*).²⁸ When our pipeline is tested on a larger consistently annotated data set and adapted to the removal of false positives due to vessel shadowing, it could be used as an analysis tool in an OCT device. Many of today's device algorithms (eg, for retinal thickness segmentation) can show errors in severe pathologic cases and clinicians' feedback corrections of the device software are needed. Owing to our fast learning time, we can even envision to improve our algorithm by using segmentation corrections of clinicians for new learning sets centrally at the device manufacturer further improving by its use in daily clinical life.

In comparison with other automated retinal fluid quantification attempts in the literature, our pipeline's learning step is exceptionally efficient due to dimension reduction. In Chen *et al*,⁵ *k*-nearest neighbor learning with ≈ 50 dimensional feature vectors is combined with graph cuts for voxel-based fluid segmentation. Our dimension reduction methodology combined with learning is pixel based and requires only two features. In contrast to Xu *et al*,⁶ we do not require a full retinal layer segmentation and reasonably smooth fluid boundaries are guaranteed by thresholding our specific pre-segmentation. As opposed to Zheng *et al*,⁴ we do not require any human interaction during the analysis procedure. Cystoid macular edema has been identified in Wilkins *et al*⁷ by an automated segmentation based on bilateral filtering and thresholding. The bilateral filtering shares similarities to the Rudin–Osher–Fatemi model, but we do not directly threshold; instead, we make a detour by computing a reasonable pre-segmentation. The algorithm in Wilkins *et al*⁷ uses orders of magnitude more training data and is more empirically tailored to their specific data set. Their test

data are few manually selected cross-sections and the results were evaluated afterwards manually, so that there seems a chance that the expert annotations are not fully unbiased from the computational results. Our pipeline works well on two independent data sets with different retinal diseases and we just use about 2 and 5% of the data sets for training, respectively. Any potential bias in graders' annotations from our computational results is avoided in the first place as the pipeline development was started after the annotations were performed. In addition, the VRC, as an independent and unbiased institution for the standardized evaluation of many multicenter clinical trials (including drug licensing trials) is used to blinding of image graders and supervisors to the purpose of their evaluation when needed.

Deep convolutional NNs have been successfully applied to detect subretinal and intraretinal cystoidal fluid.³ However, deep learning requires large amounts of training data. In Lee *et al*,²⁹ central OCT scans are linked to electronic medical records enabling the application of deep convolutional NNs. There, 80 839 training samples are used to distinguish between normal and AMD scans in 20 163 validation samples. Beyond classification, our target is fluid quantification. As the amount of accurate pixel-wise fluid annotations is often limited, we aimed to learn efficiently and use only 100 training and 4764 validation samples in AMD patients.

In conclusion, we have demonstrated that our image analysis pipeline yields good automated retinal fluid quantification in OCT images and our visual results could be useful to expert graders within their current grading processes. The pipeline is based on dimension reduction techniques with a fast learning part and low memory needs. Furthermore, it requires exceptionally few training data and is sufficiently versatile for further adjustments in the sense that current steps allow for modifications and additional steps can easily be inserted. The successful interplay of enormous dimension reduction with relatively elementary learning techniques is remarkable.

We envision several aspects for future work; for instance, we will incorporate artifact removal from vessel shadows by integrating ideas from Girard *et al*^{26,27} and Wu *et al*.²⁷ Moreover, so far, graders at the VRC annotated each OCT cross-section separately, dismissing the full spatial correlations of the 3D volumes. All of our pipeline steps are either already available or can be implemented in higher dimensions, so that we plan to perform segmentation in 3D directly. However, suitable 3D annotations are needed (*cf.* Figure 4a), but are not easy to consistently annotate

even for an experienced grader or retinal specialist. Both, vessel shadow removal and 3D implementations are work in progress.

Summary

What was known before

- Automated OCT image analysis is the future in ophthalmology.
- Fluid regions are regions of interest for macular OCT evaluation.
- Many algorithms published until now use large computational time and are not feasible in clinical routine.

What this study adds

- Validated algorithm for retinal fluid identification in the two main retinal diseases (neovascular AMD and DME).
- Comparison with standardized reading-center controlled manually annotated OCT data.
- Fast learning part allows use in clinical routine.

Conflict of interest

The authors declare no conflict of interest.

Acknowledgements

This work was funded by the Vienna Science and Technology Fund (WWTF) through project VRG12-009.

References

- 1 Schmidt-Erfurth U, Klmscha S, Waldstein SM, Bogunovic H. A view of the current and future role of optical coherence tomography in the management of age-related macular degeneration. *Eye* 2017; **31**: 26–44.
- 2 LeCun Y, Bengio Y, Hinton G. Deep learning. *Nature* 2015; **521**: 436–444.
- 3 Schlegl T, Glodan A, Podkowinski D, Waldstein SM, Gerendas BS, Schmidt-Erfurth U *et al*. Automatic segmentation and classification of intraretinal cystoid fluid and subretinal fluid in 3D-OCT using convolutional neural networks. *Invest Ophthalmol Vis Sci* 2015; **560**(7).
- 4 Zheng Y, Sahni J, Campa C, Stangos AN, Raj A, Harding SP. Computerized assessment of intraretinal and subretinal fluid regions in spectral-domain optical coherence tomography images of the retina. *Am J Ophthalmol* 2013; **155**: 277–286.
- 5 Chen X, Niemeijer M, Zhang L, Lee K, Abramoff MD, Sonka M. 3D segmentation of fluid-associated abnormalities in retinal OCT: probability constrained graph-search—graph-cut. *IEEE Trans Med Imaging* 2012; **31**(8): 1521–1531.
- 6 Xu X, Lee K, Zhang L, Sonka M, Abramoff MD. Stratified sampling voxel classification for segmentation of intraretinal and subretinal fluid in longitudinal clinical OCT data. *IEEE Trans Med Imaging* 2015; **340**(7): 1616–1623.
- 7 Wilkins G, Houghton O, Oldenburg A. Automated segmentation of intraretinal cystoid fluid in optical

- coherence tomography. *IEEE Trans Biomed Eng* 2012; **59**(4): 1109–1114.
- 8 Grady L. Random walks for image segmentation. *IEEE Trans Pattern Anal Mach Intell* 2006; **28**(11): 1768–1783.
- 9 Czaja W, Ehler M. Schrödinger Eigenmaps for the analysis of biomedical data. *IEEE Trans Pattern Anal Mach Intell* 2013; **35**(5): 1274–1280.
- 10 Schölkopf B, Smola A. *Learning with Kernels*. MIT Press: Cambridge, MA, 2002.
- 11 Roweis ST, Saul LK. Nonlinear dimensionality reduction by locally linear embedding. *Science* 2000; **290**(5500): 2323–2326.
- 12 Daugman J. Complete discrete 2-D Gabor transforms by neural networks for image analysis and compression. *IEEE Trans Acoust* 1988; **36**(7): 1169–1179.
- 13 Rudin L, Osher S, Fatemi E. Nonlinear total variation based noise removal algorithm. *Phys D* 1992; **60**: 259–268.
- 14 Storath M, Weinmann A. Fast partitioning of vector-valued images. *SIAM J Imaging Sci* 2014; **7**: 1826–1852.
- 15 Soares JVB, Leandro JG, Cesar RM, Jelinek HF, Cree MJ. Retinal vessel segmentation using the 2-D Gabor wavelet and supervised classification. *IEEE Trans on Med Imaging* 2006; **25**(9): 1214–1222.
- 16 Mumford D, Shah J. Optimal approximations by piecewise smooth functions and associated variational problems. *Comm Pure Appl Math* 1989; **42**(5): 577–685.
- 17 Potts R. Some generalized order-disorder transformations. *Proc Cambridge Philos Soc* 1952; **48**: 106–109.
- 18 Bishop CM. *Pattern Recognition and Machine Learning*. Springer-Verlag: New York, 2006.
- 19 Witten IH, Frank E. *Data Mining: Practical Machine Learning Tools and Techniques*. Elsevier, 2005.
- 20 Garvin MK, Abramoff MD, Wu X, Russell SR, Burns TL, Sonka M. Automated 3-D intraretinal layer segmentation of macular spectral-domain optical coherence tomography images. *IEEE Trans Med Imaging* 2009; **28**(9): 1436–1447.
- 21 Haeker M, Abramoff M, Kardon R, Sonka M. Segmentation of the surfaces of the retinal layer from OCT images. *Med Image Comput Comput Assist Interv* 2006; **9**(Pt 1): 800–807.
- 22 Waldstein SM, Philip A-M, Leitner R, Simader C, Langs G, Gerendas BS et al. Correlation of 3-dimensionally quantified intraretinal and subretinal fluid with visual acuity in neovascular age-related macular degeneration. *JAMA Ophthalmol* 2016; **134**(2): 182–190.
- 23 Varnousfaderani ES, Wu J, Vogl W-D, Philip A-M, Montuoro A, Leitner R et al. A novel benchmark model for intelligent annotation of spectral-domain optical coherence tomography scans using the example of cyst annotation. *Comput Methods Programs Biomed* 2016; **130**: 93–105.
- 24 Estes NAM. Computerized interpretation of ECGs. *Circ Arrhythm Electrophysiol* 2013; **6**: 2–4.
- 25 Wetherell H. Can we trust our ECG machines? *Br J Cardiol* 2014; **21**: 62–63.
- 26 Girard MJA, Strouthidis NG, Ethier CR, Mari JM. Shadow removal and contrast enhancement in optical coherence tomography images of the human optic nerve head. *Invest Ophthalmol Vis Sci* 2011; **52**(10): 7738–7748.
- 27 Wu J, Gerendas BS, Waldstein SM, Simader C, Schmidt-Erfurth U. Automated vessel shadow segmentation of fovea-centered spectral-domain images from multiple OCT devices. *Proc. SPIE 9034, Medical Imaging* 2014; **9034**: 903403.
- 28 Ehler M, Dobrosotskaya J, Cunningham D, Wong WT, Chew EY, Czaja W et al. Modeling photo-bleaching kinetics to create high resolution maps of rod rhodopsin in the human retina. *PLoS One* 2015; **10**(7): e0131881.
- 29 Lee CS, Baughman DM, Lee AY. Deep learning is effective for classifying normal versus age-related macular degeneration optical coherence tomography images. *Ophthalmol Retina*, in press 2017.

Supplementary Information accompanies this paper on Eye website (<http://www.nature.com/eye>)

Supplementary Material: Supervised learning and dimension reduction techniques for quantification of retinal fluid in optical coherence tomography images

1. LOCAL FEATURES

1.1. **Gabor filter.** For $x \in \mathbb{R}^2$, a Gabor filter is given by

$$(1) \quad h_{\sigma, \alpha, \lambda}(x) = g(D_{\sigma} R_{\alpha} x) e^{-i\lambda \langle (1,0)^{\top}, R_{\alpha} x \rangle},$$

where $g(x) = e^{-\|x\|^2}$ is the Gaussian window function, λ is a frequency (wavelength) parameter, and α , σ correspond to rotation (orientation) and anisotropic dilation with the rotation and dilation matrices

$$R_{\alpha} = \begin{pmatrix} \cos(\alpha) & -\sin(\alpha) \\ \sin(\alpha) & \cos(\alpha) \end{pmatrix}, \quad D_{\sigma} = \begin{pmatrix} \sigma^{\frac{1}{2}} & 0 \\ 0 & \sigma^{-\frac{1}{2}} \end{pmatrix},$$

respectively. The 2D Gabor filter can be viewed as a Gaussian window modulated by a complex sinusoidal plane wave of particular wavelengths and orientations. By choosing different wavelengths and orientations one creates a filter bank, in which each channel relates to different textures of the image. Here, we choose the maximum across all channels, while using 3 wavelengths λ and 5 orientations α with the scaling $\sigma = 1/\lambda$ yielding one single feature image. For further mathematical background on Gabor and time-frequency analysis, we refer to [1], [2].

1.2. **Rudin-Osher-Fatemi variational energy minimization** [3]. For some initial image $z : \Omega \rightarrow \mathbb{R}$ on a continuous region $\Omega \subset \mathbb{R}^2$ and $\gamma > 0$, we aim for

$$(2) \quad \arg \min_u \gamma \cdot \|Du\|_{L^1} + \int_{\Omega} (u(x) - z(x))^2 dx,$$

where the minimization is over all functions u with bounded variation and Du denotes the weak derivative. The second term provides approximation to the original data and the first term enforces smoothness between interfaces. The parameter γ controls the balance between the two penalties. To actually process images, a suitable discretization is necessary. Indeed, for a given discrete $n \times m$ image z with values in $\mathbb{R}$, the discrete domain version is given by

$$(3) \quad \arg \min_{u \in \mathbb{R}^{n \times m}} \sum_{ij} \frac{1}{2} (u_{ij} - z_{ij})^2 + \gamma \sqrt{(\delta_x^+ u)_{ij}^2 + (\delta_y^+ u)_{ij}^2 + \beta^2},$$

with $\beta > 0$ small and discrete differential operators $(\delta_x^+ u)_{ij} := (u_{(i+1)j} - u_{ij})$ and $(\delta_y^+ u)_{ij} := (u_{i(j+1)} - u_{ij})$, see [4] for further details and a suitable fixed point iteration for the actual minimization.

2. DIMENSION REDUCTION I: PRE-SEGMENTATION

For some initial image $z : \Omega \rightarrow \mathbb{R}^r$ on a continuous region $\Omega \subset \mathbb{R}^2$ and $\gamma > 0$, find the piecewise constant minimizer of

$$(4) \quad \arg \min_u \gamma \cdot \|\nabla u\|_0 + \int_{\Omega} (u(x) - z(x))^2 dx,$$

where $u : \Omega \rightarrow \mathbb{R}^r$ denotes a bounded piecewise constant function and $\|\nabla u\|_0$ the total boundary length of its partitioning, cf. [5], [6], [7] for details. The second term provides approximation to the original data and the first term enforces the partitioning. Again the parameter γ controls the balance between the two penalties. In [7] a fast strategy for solving

a discretization of (3) is presented. For a given discrete $n \times m$ image z with values in $\mathbb{R}^r$ the discrete domain version is given by

$$(5) \quad \operatorname{argmin}_{u \in \mathbb{R}^{m \times n \times r}} \gamma \cdot \sum_{i,j} \sum_{(a,b) \in N} \omega_{ab} \cdot [u_{i,j,:} \neq u_{i+a,j+b,:}] + \sum_{i,j,k} |u_{ijk} - z_{ijk}|^2,$$

where $u_{i,j,:}$ denotes the label vector of u at the pixel (i, j) , N denotes a neighborhood system with nonnegative weights ω and the Iverson bracket $[\cdot]$ equals one if the expression in the bracket is true and zero otherwise.

3. BLAND-ALTMAN PLOT

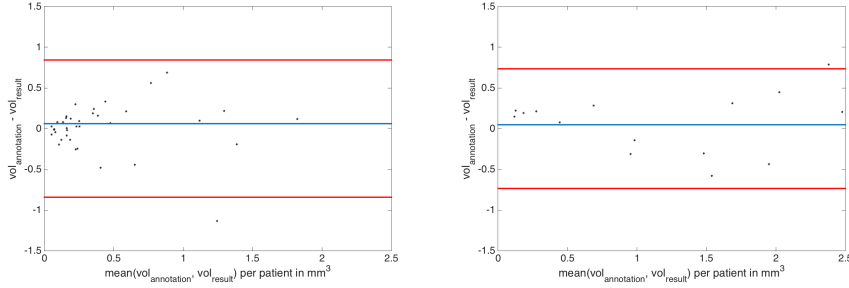


FIGURE 1

Bland-Altman plots of the volumes for the NN quantification results (left) AMD, (right) DME with lines showing the bias (blue) with ± 2 standard deviations (red).

REFERENCES

- [1] H.G. Feichtinger and T. Strohmer. *Gabor Analysis and Algorithms - Theory and Applications*. Springer Science + Business Media New York, 1998.
- [2] K. Gröchenig. *Foundations of Time-Frequency Analysis*. Springer Science + Business Media, LLC, 2001.
- [3] L. Rudin, S. Osher, and E. Fatemi. Nonlinear total variation based noise removal algorithm. *Physica D*, 60(259-268), 1992.
- [4] A. Breger. Mathematische Bildverarbeitung mit Funktionen beschränkter Variation. Bachelor thesis, 2013.
- [5] D. Mumford and J. Shah. Optimal approximations by piecewise smooth functions and associated variational problems. *Communications on Pure and Applied Mathematics*, XLII:577 – 685, 1989.
- [6] Y. Boykov, O. Veksler, and R. Zabih. Fast approximate energy minimization via graph cuts. *Cornell University*, 2014.
- [7] M. Storath and A. Weinmann. Fast partitioning of vector-valued image. *SIAM J. Imaging sciences*, 7:1826–1852, 2014.

(P4)

An Amplified-Target Loss Approach for Photoreceptor Layer Segmentation in Pathological OCT Scans

Conference proceeding: Ophthalmic Medical Image Analysis (OMIA), Lecture Notes in Computer Science, vol 11855, Springer

Authors: A. Breger, J. I. Orlando, H. Bogunovic, S. Riedl, B. S. Gerasdas, M. Ehler, U. Schmidt-Erfurth

Published: 2019

Links: https://link.springer.com/chapter/10.1007/978-3-030-32956-3_4
arXiv: <https://arxiv.org/abs/1908.00764>

DOI: 10.1007/978-3-030-32956-3

Contribution: Conceptual idea; Design of the numerical experiments; Writing.

An amplified-target loss approach for photoreceptor layer segmentation in pathological OCT scans

Anna Breger^{1*}, José Ignacio Orlando^{2*}, Hrvoje Bogunović², Sophie Riedl²,
Bianca S. Gerendas², Martin Ehler¹, and Ursula Schmidt-Erfurth²

¹ Department of Mathematics, University of Vienna, Austria

² Christian Doppler Laboratory for Ophthalmic Image Analysis (OPTIMA),
Department of Ophthalmology and Optometry, Medical University of Vienna, Austria

Abstract. Segmenting anatomical structures such as the photoreceptor layer in retinal optical coherence tomography (OCT) scans is challenging in pathological scenarios. Supervised deep learning models trained with standard loss functions are usually able to characterize only the most common disease appearance from a training set, resulting in sub-optimal performance and poor generalization when dealing with unseen lesions. In this paper we propose to overcome this limitation by means of an augmented target loss function framework. We introduce a novel amplified-target loss that explicitly penalizes errors within the central area of the input images, based on the observation that most of the challenging disease appearance is usually located in this area. We experimentally validated our approach using a data set with OCT scans of patients with macular diseases. We observe increased performance compared to the models that use only the standard losses. Our proposed loss function strongly supports the segmentation model to better distinguish photoreceptors in highly pathological scenarios.

1 Introduction

Supervised deep learning techniques have revolutionized the field of medical image segmentation [1], particularly with fully convolutional neural network architectures such as the U-Net [2]. To learn these networks, a loss function L is optimized using gradient based approaches and backpropagation. This function is usually defined in terms of metrics that quantify the discrepancies between a trustworthy/ground truth labelling and the predicted segmentation.

In this typical framework a loss function is not explicitly tailored to aim for a specific feature in the target space. Hence, the network firstly learns the dominating characteristics of the target images in the training set, and its remaining capacity is gradually devoted to characterize other less prevalent target features. This becomes an issue when dealing with highly pathological data, where lesions or disease appearance might significantly differ between patients. To overcome

* equal contribution

this limitation, some authors proposed to train segmentation models using a linear combination of different losses such as cross-entropy and Dice [3]. However, these metrics are still computed from the same target representation, so they do not enhance a specific target feature. In this paper we propose to extend this idea by using the framework of *augmented target loss functions*, introduced in [4]. Rather than relying on a single or a linear combination of loss functions defined on the original prediction and target space, Breger *et al.* [4] proposed to compute the loss on alternative representations of the predictions and targets, obtained by applying differentiable transformations T that enhance specific characteristics.

This paper focuses on the application of an augmented target loss function for photoreceptor layer segmentation in retinal optical coherence tomography (OCT) scans of patients with macular diseases. OCT is the state-of-the-art technique for imaging the retina, as it brings volumetric information through a stack of 2D scans (B-scans) at a micrometric resolution [5]. Ophthalmic disorders such as diabetic macular edema (DME), retinal vein occlusion (RVO) and age-related macular degeneration (AMD) gradually affect photoreceptors while progressing. The abnormal accumulation of fluid due to these diseases significantly alters the retina, eventually leading to photoreceptor cell death. This last characteristic can be noticed through OCT imaging: first as a pathological thinning of the photoreceptor layer, and more lately as complete disruptions on it (Fig. 1, right). It has been observed that these abnormalities are highly correlated with focal vision impairment [6] and visual acuity loss when located at the central area of the retina [7]. Hence, the automated characterization of the morphology of the photoreceptor layer is relevant for efficient quantification of functional loss.

In this paper we build on top of the architectural innovations proposed in [8] by training such a model using an augmented target loss function. Fitting the framework we introduce a novel amplified-target loss that induces further penalization to errors within the central area of the B-scans. As the most challenging pathologies are usually observed at the central area of fovea-centered OCT scans, our hypothesis is that incorporating this loss function as a kind of regularizer enforces the network to better characterize disease appearance. We validate our approach using a series of OCT scans of patients with AMD, DME and RVO. Our results empirically show that the proposed loss functions improve the performance within the central millimeters of the retina compared to using traditional losses without compromising the performance in the entire volume.

2 Methods

2.1 Augmented target loss functions for image segmentation

In a supervised learning problem we aim to learn a function f with $f_{\theta}(x) \approx y$, where θ denotes the free parameters and $S = \{(x, y)^{(i)}\}, 1 < i < N$ is a given training set with pairs of inputs x and ground truth labels y . In the context of image segmentation, x corresponds to an input image, y and $\hat{y}$ are manual and predicted segmentations and f_{θ} is some segmentation model (e.g. a fully convolutional neural network such as the U-Net [2]).

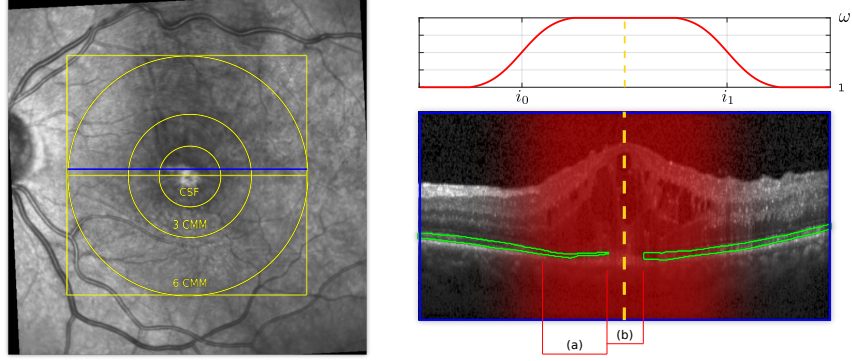


Fig. 1. Left: scanning laser ophthalmoscopy (SLO) of a patient with RVO. The square indicates the area captured by the OCT volume and the rings represent the central subfield (CSF) and the 3 and 6 central millimeters (3 CMM and 6 CMM). The blue line highlights the B-scan showed in the right side. Right: CSF B-scan with photoreceptor layer annotation (green) with (a) disruptions and (b) abnormal thinning. The red heat map represents the weighting strategy applied in our loss function. The central coordinate of the image is indicated with the yellow dotted line, and a profile of the weighting strategy is illustrated on top of the B-scan.

To adjust the weights θ from the chosen network structure f_θ , a loss function L is minimized using gradient based optimization. L is a piecewise differentiable loss function, e.g. cross-entropy (CE) or mean squared error (MSE), that measures the pixel-wise differences between $\hat{y}$ and y . In standard settings no specific areas of the images are penalized more than others. Thus, the parameters θ are mostly adjusted to characterize those features from the training set that have the most impact on the overall loss. Although this might be helpful to segment healthy anatomy, in pathological scenarios the network will overfit the prevalent features unless explicit regularization is imposed during training.

Here, we propose to use the framework of augmented target (AT) loss functions, introduced in [4]. These losses take into account prior knowledge of target characteristics via error estimation in transformed target spaces. The framework can be applied to any supervised learning problem based on loss optimization where additional information about the target data is available, provided it can be formulated as a transformation function T . The transformation may correspond to any piecewise differentiable function on the target space that yields a more beneficial representation of some known target characteristic.

Following [4], the AT loss functions L_{AT} is a linear combination of losses applied to transformed targets. Its general form is:

$$L_{AT} = \sum_{j=1}^d \lambda_j \cdot L^j(\{T_j(y_i)\}_i, \{T_j(\hat{y}_i)\}_i), \quad (1)$$

where $\lambda_j > 0$ corresponds to some weight, T_j to a specific transformation and L^j to some loss function, for all $j \in \{1, \dots, d\}$.

Setting typically T_1 to the identity and L^1 to a standard loss, the additional terms in the L_{AT} loss act as amplified target information, yielding a new optimization problem:

$$\hat{\theta} = \arg \min_{\theta} \{ \lambda_1 \cdot L_1(\{y_i\}_i, \{\hat{y}_i\}_i) + \sum_{j=2}^d \lambda_j \cdot L^j(\{T_j(y_i)\}_i, \{T_j(\hat{y}_i)\}_i) \}, \quad (2)$$

where the weights λ_1 and $\{\lambda_j\}_{j=2}^d$ control the balance between the main loss and the regularization terms respectively.

2.2 Amplified-target loss functions for photoreceptor layer segmentation

We experimentally study the AT loss function framework in the context of photoreceptor layer segmentation in pathological OCT scans. We tailor a so called *amplified-target loss* in which a transformation T is designed to bring an increased penalty to errors within the central area of the images. This loss is intended to incorporate the prior knowledge that abnormalities such as pathological thinnings and disruptions of the photoreceptor layer are more common in the central millimeters of the foveal area. To do so, we define a transformation $T(y_i) = W \odot y_i$, where $\odot$ denotes the Hadamard (entrywise) product, y_i corresponds to the given binary targets and W represents a weighting matrix that encodes a penalization weight for errors. This operation can analogously be applied to the predictions $\hat{y}_i$. Fig. 1 graphically illustrates the design of the weighting matrix W . Formally, we define $W = G_{\sigma} * V$, where G_{σ} stands for a Gaussian filter with standard deviation σ . We define V as:

$$V_{j,i} := \begin{cases} \omega & \text{for } i_0 < i < i_1 \text{ and all } j, \\ 1 & \text{otherwise,} \end{cases} \quad (3)$$

where ω denotes the maximum weight assigned to the central area and $[i_0, i_1]$ is the horizontal interval of the image that is amplified. The Gaussian filter G_{σ} is used to smooth the penalization factor within the edges of the interval.

Following the formulation in (2), we can then redefine our empirical risk minimization problem as

$$\hat{\theta} = \arg \min_{\theta} \{ \lambda_1 \cdot L^1(\{y_i\}_i, \{\hat{y}_i\}_i) + \lambda_2 \cdot L^2(\{W \odot y_i\}_i, \{W \odot \hat{y}_i\}_i) \}, \quad (4)$$

where we choose $\lambda_1, \lambda_2 \in \mathbb{R}$ and $L^1 = L^2$ as CE or MSE losses.

3 Experimental setup

3.1 Materials

Our method was trained and tested on an in-house data set with 53 Spectralis OCT volumes of patients suffering from DME (16), RVO (27) and AMD (10).

Each image comprises 496×512 pixels per B-scan, 49 B-scans per volume. All the B-scans were manually annotated by certified readers under the supervision of a retina expert, who modified the labels when necessary to obtain ground truth segmentations. The set was randomly divided into a training, a validation and a test set, each of them with 34, 4 and 15 scans, respectively, with approximately the same distribution of diseases and percentages of disrupted columns per B-scan (or A-scans).

3.2 Network architecture and training setup

We used the photoreceptor segmentation network described in [8] in our experiments (note that any other architecture could be applied within our framework). We used as baselines CE and MSE comparing it to the adapted AT loss.

Every configuration was trained at a B-scan level with a batch size of 2 images, using Adam optimization and early stopping. Hence, training was stopped if the validation loss did not improve for the last 45 epochs. The learning rate was set to $\eta = 0.0001$, and divided by 2 if the validation loss was not improved during the last 15 epochs. Data augmentation was used in the form of random horizontal flippings. Binary segmentations were retrieved as in [8] by thresholding the softmax scores of the photoreceptors class using the Otsu algorithm.

4 Results and Discussion

We evaluated the performance for segmenting the photoreceptor layer using the volume-wise Dice index, at the CSF, the 3 CMM, the 3-1 ring and the full volume (Fig. 1, left). All the experiments with our AT loss functions were conducted using fixed values for $\sigma = \frac{1}{16}X$, $i_0 = \frac{1}{4}X$ and $i_1 = \frac{3}{4}X$ (with $X = 512$ being the horizontal size of the B-scans, in pixels), without optimizing them on the validation set. Different configurations for $\omega = 2^k$, $k \in \{1, \dots, 5\}$ and λ_1 and $\lambda_2 \in \{0.001, 0.01, 0.1, 1, 2, 4, 8\}$ were analyzed, and the best configuration according to Dice index on the validation set was then fixed to allow a fair comparison on the test set. From this model selection step, we observed that $\omega = 8$, $\lambda_1 = 1$ and $\lambda_2 = 8$ reported the best performance for the AT loss with categorical cross-entropy (CE), and $\omega = 32$, $\lambda_1 = \lambda_2 = 1$ for the AT loss with mean square error (MSE).

Fig. 2 depicts boxplots with the quantitative performance of each model on the test set, compared with their corresponding baselines trained only with CE and MSE, for each evaluation area. The mean and standard deviation values of the Dice index are presented in Table 1. The incorporation of the AT loss allows to perform consistently better in all the cases, with the best results reported by the MSE loss. Statistical analysis using one-tail Wilcoxon sign-rank tests at a significance level $\alpha = 0.05$ showed that the model trained with MSE + AT loss reported significantly higher Dice values in the CSF area compared to using CE + AT loss or only MSE ($p < 0.0171$). These differences were not statistical significant with respect to the model trained with CE ($p = 0.1902$).

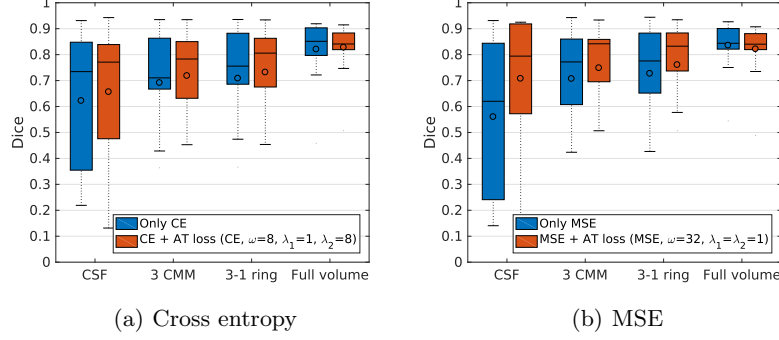


Fig. 2. Volume-wise Dice values for all the evaluated models and our proposed approach in each evaluation area. Circles indicate mean values. CSF: central subfield (1 central millimeter). 3 CMM: three central millimeter. 3-1 ring: area between CSF and 3 CMM.

Table 1. Volume-wise mean $\pm$ standard deviation Dice values in the test set for each photoreceptor segmentation model and the different areas.

Method	CSF	3 CMM	3-1 ring	Full volume
CE loss	0.622 $\pm$ 0.271	0.691 $\pm$ 0.242	0.708 $\pm$ 0.242	0.820 $\pm$ 0.118
CE + AT loss (CE, $\omega = 8$, $\lambda_1 = 1$, $\lambda_2 = 8$)	0.656 $\pm$ 0.256	0.718 $\pm$ 0.218	0.732 $\pm$ 0.218	0.828 $\pm$ 0.100
MSE loss	0.560 $\pm$ 0.303	0.707 $\pm$ 0.223	0.727 $\pm$ 0.223	0.835 $\pm$ 0.096
MSE + AT loss (MSE, $\omega = 32$, $\lambda_1 = \lambda_2 = 1$)	0.708 $\pm$ 0.254	0.749 $\pm$ 0.215	0.760 $\pm$ 0.213	0.821 $\pm$ 0.102

When comparing the Dice values at the 3-1 ring, the MSE with AT loss model reported statistically significant better results than using only CE or MSE ($p < 0.0042$), which is consistent with its behavior in the 3 CMM ($p < 0.0416$). No statistically significant differences in performance were observed at the full volume level (two-tails test, $p > 0.0730$).

We qualitatively analyzed the segmentation and score maps using the CE and MSE combined with the AT loss. Fig. 3 depicts segmentation results in a central B-scan from the test set, with score maps represented as heatmaps. Using MSE produces noisy scores within the lateral areas of the B-scans, and therefore spurious elements in the segmentation. CE, on the contrary, results in smoother score maps, although with few false negatives in the vicinity of subretinal fluid. This behavior is linked to the one observed in Table 1, where the MSE + AT loss model reported higher Dice in the central area than using CE, and smaller values in the full volume. The model trained with only MSE performs poorly in the CSF, the 3 CMM and the 3-1 ring, which indicate that it struggles to deal with pathologies. Similarly, the high performance at a volume level indicates that

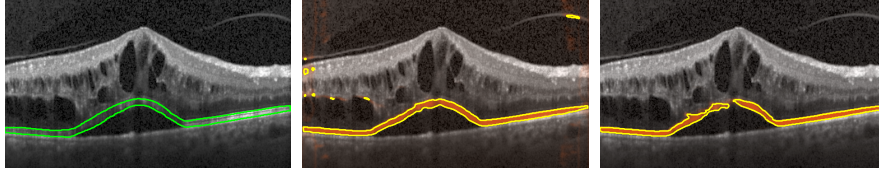


Fig. 3. Qualitative effect of the loss selection in the pixel score values. From left to right: manual annotation (green), score map (orange) and binary segmentation (yellow) obtained with MSE + AT loss (MSE, $\omega = 32$, $\lambda_1 = \lambda_2 = 1$) and CE + AT loss (CE, $\omega = 8$, $\lambda_1 = 1$, $\lambda_2 = 8$).

it can better characterize normal appearances. When using MSE + AT loss, a significant reduction in the amount of false negatives occurs at the central areas. However, as mentioned before, the score maps are noisy at the borders of the B-scans, which causes a drop in the full volume Dice. The model trained with CE + AT loss is less accurate at the center than the one trained with MSE + AT loss, but it still outperforms the baseline approaches. Moreover, at a volume basis the CE + AT loss remains competitive with respect to the one trained only with CE loss.

Finally, Fig. 4 presents qualitative results in exemplary central B-scans from our test set obtained both by the models trained with CE only and with CE + AT loss. Our approach produced more anatomically plausible segmentations than the standard CE loss in pathological areas with subretinal fluid (Fig. 4 (a) and (b)) or large disruptions (Fig. 4 (c)).

5 Conclusions

In this paper we proposed to use the framework of augmented target loss functions for photoreceptor layer segmentation in pathological OCT scans. We define an amplified-target loss incorporating a transformation that weights the central area of the input B-scans to further penalize errors committed in this region. We experimentally observed that this straightforward approach allows to significantly improve performance within the central millimeters of fovea-centered OCT scans, without affecting the overall performance in the entire volume. These results indicate that the proposed AT loss function acts as a form of regularization, better characterizing photoreceptors appearance within highly pathological regions. We are currently exploring new alternatives to identify the regions to weight and to learn their corresponding weights. Further experiments are also performed to evaluate our approach in the context of other OCT based applications such as fluid segmentation and using OCT scans from other vendors.

Acknowledgements

This work is funded by WWTF AugUniWien/FA7464A0249 (MedUniWien); VRG12-009 (UniWien). We thank NVIDIA Corporation for donating a GPU.

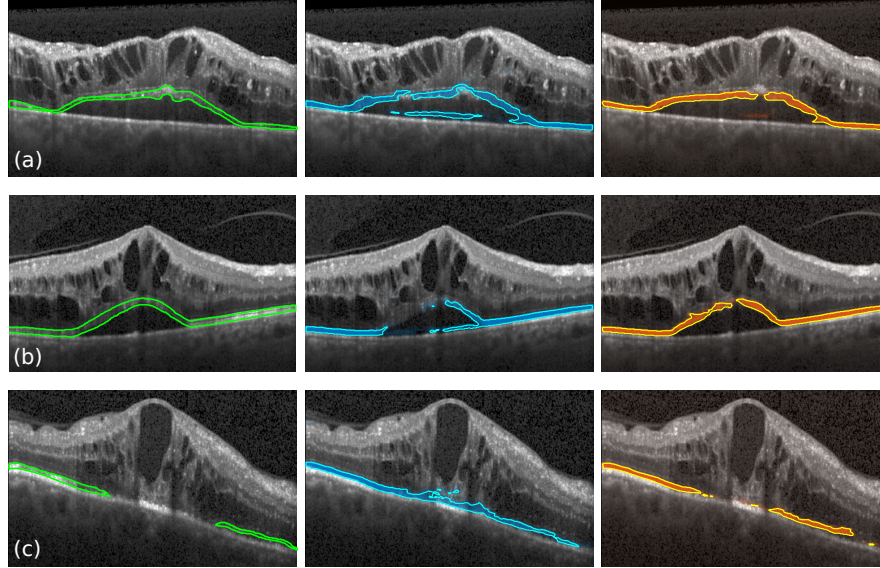


Fig. 4. Qualitative results in central B-scans from the test set. From left to right: manual annotations (green), results with only CE loss (blue) and results with CE + AT loss (CE, $\omega = 8$, $\lambda_1 = 1$, $\lambda_2 = 8$).

References

1. Litjens, G., et al.: A survey on deep learning in medical image analysis. *Medical image analysis* **42** (2017) 60–88
2. Ronneberger, O., Fischer, P., Brox, T.: U-Net: Convolutional Networks for Biomedical Image Segmentation. In Navab, N., Hornegger, J., Wells, W.M., Frangi, A.F., eds.: MICCAI 2015, Springer (2015) 234–241
3. Khened, M., Kollerathu, V.A., Krishnamurthi, G.: Fully convolutional multi-scale residual densenets for cardiac segmentation and automated cardiac diagnosis using ensemble of classifiers. *Medical image analysis* **51** (2019) 21–45
4. Breger, A., Orlando, J.I., Harár, P., Doerfler, M., Klimscha, S., Grechenig, C., Gerendas, B., Schmidt-Erfurth, U., Ehler, M.: On orthogonal projections for dimension reduction and applications in augmented target loss functions for learning problems. *Journal of Mathematical Imaging and Vision (JMIV)*, Springer (2019)
5. Schmidt-Erfurth, U., Sadeghipour, A., Gerendas, B.S., Waldstein, S.M., Bogunović, H.: Artificial intelligence in retina. *Progress in Retinal and Eye Research* (aug 2018)
6. Takahashi, A., et al.: Photoreceptor damage and reduction of retinal sensitivity surrounding geographic atrophy in age-related macular degeneration. *American journal of ophthalmology* **168** (2016) 260–268
7. Gerendas, B.S., et al.: OCT biomarkers predictive for visual acuity in patients with diabetic macular edema. *IOVS* **58**(8) (2017) 2026–2026
8. Orlando, J.I., et al.: U2-Net: A Bayesian U-Net model with epistemic uncertainty feedback for photoreceptor layer segmentation in pathological OCT scans. In: ISBI. (2019)

(P5)

On the Reconstruction Accuracy of Multi-Coil MRI with Orthogonal Projections

Status: preprint (2019)

Authors: A. Breger, G. Ramos Llorden, G. Vegas Sanchez - Ferrero, W. S. Hoge, M. Ehler, and C-F. Westin

Link: <https://arxiv.org/abs/1910.13422>

Contribution: Conceptual idea; Theory development; Design and performance of numerical experiments; Main writing.

On the reconstruction accuracy of multi-coil MRI with orthogonal projections

Anna Breger¹, Gabriel Ramos Llorden², Gonzalo Vegas Sanchez - Ferrero³, W. Scott Hoge³, Martin Ehler¹, and Carl-Fredrik Westin⁴

¹Department of Mathematics, University of Vienna, Austria

²Department of Psychiatry, Brigham and Women's Hospital, Harvard Medical School,
Boston, MA

³Department of Radiology, Brigham and Women's Hospital, Boston, MA

⁴Laboratory of Mathematics in Imaging, Brigham and Women's Hospital, Harvard Medical
School, Boston, MA

September 2019

Abstract

MRI signal acquisition with multiple coils in a phased array is nowadays commonplace. The use of multiple receiver coils increases the signal-to-noise ratio (SNR) and enables accelerated parallel imaging methods. Some of these methods, like GRAPPA or SPIRiT, yield individual coil images in the k-space domain which need to be combined to form a final image. Coil combination is often the last step of the image reconstruction, where the root sum of squares (rSOS) is frequently used. This straightforward method works well for coil images with high SNR, but can yield problems in images with artifacts or low SNR in all individual coils. We aim to analyze the final coil combination step in the framework of linear compression, including principal component analysis (PCA). With two data sets, a simulated and an in-vivo, we use random projections as a representation of the whole space of orthogonal projections. This allows us to study the impact of linear compression in the image space with diverse measures of reconstruction accuracy. In particular, the L_2 error, variance, SNR, and visual results serve as performance measures to describe the final image quality. We study their relationships and observe that the L_2 error and variance strongly correlate, but as expected minimal L_2 error does not necessarily correspond to the best visual results. In terms of visual evaluation and SNR, the compression with PCA outperforms all other methods, including rSOS on the uncompressed image space data.

1 Introduction

Magnetic Resonance Imaging (MRI) is a unique medical imaging modality that provides excellent image quality without ionizing radiation, but on the other hand is relatively slow. The acquired data, sampled in the k-space domain, corresponds to the Fourier transform of the spatial-domain MR image. To reconstruct an accurate MR image, the total amount of k-space data that must be acquired to avoid reconstruction artifacts is dictated by sampling theory.

A turning point in MRI reconstruction was the implementation of phased-array coils in the acquisition pipeline. With phased-array coils, the k-space data is acquired in multiple receivers/coils simultaneously. On top of increasing the signal-to-noise (SNR) ratio, a phased-array coil acquisition allows accelerating the total acquisition time [22]. These methods, known as Parallel Imaging (PI) methods, exploit the fact that there exist complementary k-space data information in each coil/receiver. In PI, the k-space data set is undersampled, but the multiple measurements allow missing k-space samples to be recovered through an inverse problem reconstruction framework. See [14] for an overview of basic reconstruction algorithms for PI and their history.

The computational costs and required memory of the reconstruction algorithms highly depend on the dimension of the phased array, i.e. the number of receiver coils. Several coil compression methods have been developed that reduce to a smaller set of virtual channels without a significant loss of SNR, see e.g. [29],[8],[26]. Moreover, PCA-based methods have shown to even have a beneficial denoising effect, e.g. [6]. In [5], the optimal linear projection for coil compression based on the resulting SNR is derived. Note that this is related to our analysis approach, but we work in the image domain instead of the k-space domain.

After the obligatory compression step, PI reconstruction methods can be applied and generate results in different output spaces. Some of the methods operate in the image domain, e.g. SENSE [20]. Other methods operate in the k-space domain and generate a fully-sampled k-space array per each receiver/coil, e.g. GRAPPA or SPIRiT [15, 19]. To form a spatial-domain image that can be used for analysis, quantification, or visualization purposes, the k-space domain data is inversely Fourier transformed and the resulting individual coil images need to be combined into a single final image.

In [22], optimal methods to combine the arrays from phased array elements have been developed. These methods rely on detailed knowledge of the receive sensitivity of each coil. In practice, the exact position and sensitivity of each coil is not always known or not possible to estimate because of physical limitations. To overcome this issue, a method that combines the data without detailed knowledge of the coils, while preserving a high SNR, is desirable. Because of this, the root sum of squares (rSOS) method has become the standard method of combining multi-coil images in MRI [22, 25]. For arrays with high SNR, the rSOS yields nearly optimal reconstruction, whereas problems arise if all coils yield low SNR. Especially data with artifacts are problematic since rSOS weights them equally to the non-defective parts.

Noise is ever present in MRI imaging and is dominated by two sources: thermal noise in the receiver apparatus and the physiological noise from patient’s body itself. Moreover, the SNR in the coil images depends highly on the field strength, scanner hardware, and data acquisition modality (e.g. T1/T2/diffusion weighted). Extensive statistical noise analysis can be found in [1]. In [24, 18], PCA was used to jointly reconstruct multi-coil data from diffusion MRI exploiting the expected redundancy of data acquired with different gradient directions. Under the framework of random matrix theory, authors derived practical rules to accurately nullify noise-only principal components. This way, dMRI data can be effectively denoised while preserving the important information.

In this work, we aim to analyze the coil combination in a comprehensive way by comparing different performance measures. To address coil combination independent of prior knowledge, we will study the rSOS in the framework of linear compression via orthogonal projections. Random projections will serve us as a tool to study the correlation between L_2 reconstruction error and voxel variance of the varying reconstructions. Correlation analysis with random orthogonal projections has been studied in a related context in [4], yielding an underlying understanding of the relation between important information preservation features in data combination and compression.

Besides the correlation analysis, we observe that optimal L_2 reconstruction error, i.e. mean squared error (MSE), does not necessarily correspond to optimal visual evaluation of the reconstructed images. This corresponds to described problems when measuring image quality by the MSE, see e.g. [27, 13, 28]. Nevertheless, the often used peak signal-to-noise ratio (PSNR) is based on the computation of the MSE. Here, we will additionally evaluate the reconstructions regarding the SNR and visualization of the images. The results confirm an improvement by compressing image space data with PCA.

The outline is as follows: first we will explain how rSOS can be interpreted in the context of orthogonal projections. Moreover, we will describe samples of random orthogonal projections, that shall serve us as coverings enabling numerical experiments. Then we describe the two data sets, a simulated and an in-vivo, that we use for experimental investigations. The simulated data set allows us to compare the behavior with different noise levels in the coils. In the results section, we show scatter plots describing the relation between the L_2 reconstruction error and the voxel variance in reconstructed magnitude images, as well as comparing it with the SNR and visual outcomes. Finally, we will interpret and discuss the results in Section 5.

2 Methods

We aim to study the impact of linear compression in coil combination with rSOS. Such a combination step is necessary for phased-array data that has been processed with GRAPPA, SPIRiT, or other PI methods in k-space. Before combining the fully sampled, multidimensional image space data as the final step of the reconstruction pipeline, we include linear compression with orthogonal

projections.

The common reconstruction pipeline including our linear compression can be summarized as follows:

$$\hat{y}_i \in \mathbb{C}^d \xrightarrow{\text{GRAPPA}} \hat{x}_i \in \mathbb{C}^d \xrightarrow{\text{IFFT} + \text{abs}} x_i \in \mathbb{R}^d \xrightarrow{\text{projection}} px_i \in \mathbb{R}^k \xrightarrow{\text{rSOS}} \|px_i\|_2 \in \mathbb{R}$$

where $\hat{y}_i$ corresponds to an undersampled k-space voxel, $\hat{x}_i$ is fully sampled in k-space after some PI reconstruction (e.g. GRAPPA) and x_i is the image space data. Our analysis takes place on the image space data in $\mathbb{R}^d$ that is subsequently projected to $\mathbb{R}^k$ with $k < d$.

2.1 Framework of orthogonal projections

As the basis of our analysis we will use random orthogonal projections yielding different linear coil compressions. We will work with image data consisting of d channels that shall be combined into one final magnitude image. To do so, we study the space of k -dim linear subspaces of $\mathbb{R}^d$, which can be identified by orthogonal projections

$$\mathcal{G}_{k,d} = \{p \in \mathbb{R}^{d \times d} : p^2 = p, p^T = p, \text{rank}(p) = k\}, \quad (1)$$

called the Grassmannian.

Let $x = \{x_i\}_{i=1}^m \in \mathbb{R}^d$ be an image space data set with m voxels measured by d coils. Then, the final image volume is given by $\|px\|_2$ with p in $\mathcal{G}_{k,d}$, where d is the fixed number of channels and $k < d$ varies. This corresponds to projecting the d coil channels to different dimensions k and computing root sum of squares (rSOS) subsequently. Note that we can study the rSOS itself in this context, since $\text{rSOS}(x) = \|x\|_2$ and it holds for all $p \in \mathcal{G}_{k,d}$ that the expectation value

$$\mathbb{E}[\|px_i\|_2] = c \cdot \|x_i\|_2 \quad \forall x_i \in \mathbb{R}^d, \quad (2)$$

with $c = (\Gamma(\frac{k+1}{2})\Gamma(\frac{d}{2})) / (\Gamma(\frac{k}{2})\Gamma(\frac{d+1}{2}))$, where Γ denotes the Gamma function.

This is based on the Chi distribution and the fact that the length of a random unit vector projected onto a fixed k -dimensional subspace has the same distribution as the length of a unit vector in $\mathbb{R}^d$ being projected onto a random k -dimensional subspace (see e.g. [12]).

Remark 2.1 *The equality (2) allows us to analyze the rSOS itself in the framework of orthogonal projections, i.e. no added coil compression in the image space. Summing up the projected voxels obtained by a reasonably big sample set of random orthogonal projections $p = \{p_l\}_{l=1}^n \in \mathcal{G}_{k,d}$, yields approximately the rSOS combined voxel up to the constant c , i.e.*

$$\frac{1}{n} \sum_{l=1}^n \|p_l x_i\|_2 \approx c \cdot \|x_i\|_2. \quad (3)$$

Since we linearly rescale the final image volumes between $[0,1]$ to enable error estimation with the ground truth, the constant is negligible. Note that also principal component analysis (PCA) yields an orthogonal projection that lays in the

Grassmannian manifold and therefore can be studied in that context. Moreover, we will use random orthogonal projections, i.e. projections $p \in \mathcal{G}_{k,d}$ distributed according to the orthogonally invariant probability measure $\mu_{k,d}$, as samples of orthogonal projections, see e.g. [2], [3].

2.2 Measures of performance

The L_2 error will serve us as a measure of accuracy, when comparing the ground truth data with the final reconstructions, i.e. the combined image space data $x = \{x_i\}_{i=1}^m \in \mathbb{R}^d$. Note that w.l.o.g. x contains here just the voxels in the region of interest (discarding the background) and not the full image volume. For a fixed projection $p \in \mathcal{G}_{k,d}$, the error between some provided ground truth $y = \{y_i\}_{i=1}^m$ and the combined magnitude image voxels $\|px\|_2 := \{\|px_i\|_2\}_{i=1}^m$, is then given by

$$\text{Err}(y, \|px\|_2) := \frac{1}{m} \sum_{i=1}^m (y_i - \|px_i\|_2)^2. \quad (4)$$

We aim to study the relation between the reconstruction error and the variance in the reconstructed magnitude images, which can be interpreted as contrast in noise-free images. The variance of the final magnitude images depending on the projection method can be computed by

$$\text{Var}(\|px\|_2) = \frac{1}{m(m-1)} \sum_{i < j} (\|px_i\|_2 - \|px_j\|_2)^2. \quad (5)$$

Moreover, we will use the mean signal-to-noise ratio to interpret the performance. Note that in MRI it is common to use the non-squared version, i.e.

$$\text{SNR}(\|px\|_2) = \frac{1}{m} \sum_{i=1}^m \frac{\|px_i\|_2}{\sigma}, \quad (6)$$

where σ corresponds to the estimated standard deviation of the noise, see e.g. [1].

2.3 Random projections as coverings

To enable a numerical analysis, we need a finite set of orthogonal projections that represents the overall space well, i.e. covers the Grassmannian $G_{k,d}$ properly. To measure how well a set covers the underlying space, we use the definition of the covering radius.

Definition 2.2 *Let the covering radius of a finite set $\{p_1, \dots, p_n\} \subset \mathcal{G}_{k,d}$ be denoted by*

$$\rho(\{p_l\}_{l=1}^n) := \sup_{p \in \mathcal{G}_{k,d}} \min_{1 \leq l \leq n} \|p - p_l\|_F, \quad (7)$$

where $\|\cdot\|_F$ is the Frobenius norm.

Note that the smaller the covering radius, the better the finite set of projections $\{p_l\}_{l=1}^n$ represents the entire space $\mathcal{G}_{k,d}$: it yields smaller holes and the points are better spread.

Let $\mu_{k,d}$ denote the normalized Riemannian measure on $\mathcal{G}_{k,d}$ as before. According to [21], the expectation of the covering radius ρ of n random points $\{p_j\}_{j=1}^n$, independent identically distributed according to $\mu_{k,d}$, satisfies ¹

$$\mathbb{E}\rho \sim n^{-\frac{1}{k(d-k)}} \log(n)^{\frac{1}{k(d-k)}}. \quad (8)$$

Following the definition of asymptotically optimal covering in [3], the expectation yields an optimal covering radius up to a logarithmic factor $\log(n)^{\frac{1}{k(d-k)}}$. To remain flexible in the dimension of $\mathcal{G}_{k,d}$, i.e. the choice of k and d , and the number of projections n , we will work here with random projections distributed according to $\mu_{k,d}$ rather than constructing optimal covering sequences as in [3]. These random orthogonal projections can be efficiently computed by the QR decomposition of a matrix $M = QR$ with independent standard normal distributed entries [7, Theorem 2.2.2].

In the following we will use finitely many samples of random projections for the experimental analysis.

3 Data

We will run our numerical analysis on two different MR data sets: a simulated T1-weighted data set from brainweb ([10], [17], [11]) and an in-vivo data set from a head coil receiver.

3.1 Simulated data set

Simulation experiments were conducted to assess the quality in image reconstruction for different types of projections in a controlled, rigorous manner. To do so, first, a ground-truth volume was created with the popular numerical simulator BrainWeb [9]. A (magnitude) multi-slice T1-weighted volume was simulated with a Spoiled Fast Low Angle Shot (SFLASH) sequence with the following parameters: TR/TE = 20 /10 ms, flip angle of 90 degrees and ETL = 1. Matrix size: $181 \times 21 \times 76$ with isotropic voxel size of 1 mm.

Next, simulated images with 32 channels, mimicking a 32-channel coil acquisition, were created. First, synthetic coil sensitivity profiles were simulated assuming a smooth Gaussian profile [1]. Voxel-wise multiplication of those coil sensitivities by the simulated ground-truth image creates the 32 coil-based images. Uncorrelated complex Gaussian noise with zero mean and standard deviation σ was added. Finally, to simulate a magnitude-based acquisition, the absolute value of 32 noisy images were taken. The value of σ was chosen differently to recreate experiments with different noise levels.

¹We use the symbol $\sim$ to indicate that the corresponding equalities hold up to a positive constant factor on the respective right-hand side.

3.2 In-vivo data

MR data was acquired in-vivo from a healthy volunteer using a Siemens (Erlangen, Germany) 3T Prisma equipped with a 32-channel head coil receiver. An Echo-Planar diffusion sequence was employed to acquire 24 slices in a 2mm isotropic volume, with slice-thickness 2mm (TR=3.2 sec, TE=85ms, flip-angle=90, matrix size 128x128, FOV: 256mm by 256mm). A T2-weighted volume was acquired, with no diffusion weighting (a "b=0" image), followed by 62 volumes acquired using a single repeated diffusion vector and diffusion setting of b=1000. The measured EPI data was reconstructed using Dual-Polarity GRAPPA (DPG) to minimize Nyquist ghosts, see [16]. In-plane acceleration of the original data was $R = 2$. The ground-truth image was formed by motion-correcting the 62-volume diffusion-weighted series, using the Advanced Normalization Tools (ANTs) library [23], and averaging across time. The first diffusion direction was studied regarding linear coil combination and compared to the computed the ground-truth. Few outliers have been removed by using the 99.99th percentile and setting the remaining 0.01% to the maximum of the used percentile.

4 Results

Both studied image volumes consist of phased-array data with 32 channels, therefore we work with projections p in $G_{k,d}$ with $d = 32$ and varying k . In the following scatter plots we see how the L_2 reconstruction error (4) relates to the voxel variance (5) and state corresponding mean SNR values in the visualization.

Each point in the plots represents a reconstruction with linear compression and rSOS, yielding a final magnitude image volume as described in the previous sections. The symbols $+$ correspond to a projection $p \in G_{k,32}$, i.e. $\|px\|_2$, and the colors to the different dimensions k :

$$+ p \in \mathcal{G}_{1,32}, + p \in \mathcal{G}_{4,32}, + p \in \mathcal{G}_{12,32}, + p \in \mathcal{G}_{20,32}, + p \in \mathcal{G}_{28,32}, + p \in \mathcal{G}_{31,32}.$$

As described in Section 2.3 the projections p serve as a covering of the underlying space and are chosen randomly according to the orthogonally invariant probability measure $\mu_{k,d}$. For the simulated data set (see Section 3.1) $n = 500$ projections are randomly chosen for every space $\mathcal{G}_{k,32}$ with $k \in \{1, 4, 12, 20, 28, 31\}$, for the in-vivo data set (see Section 3.2) $n = 1000$.

In the scatter plots the symbol $\square$ corresponds to the rSOS, i.e. $\|x\|_2$, and $\circ$ to the compression with the orthogonal projection provided by PCA. The colors correspond to the different spaces $\mathcal{G}_{k,32}$ as stated above.

Figure 1 contains the scatter plots for 4 different noise levels in the simulated data set and Figure 2 shows corresponding image cross-sections. Figure 3 shows the scatter plot and image cross-sections of the in-vivo data set.

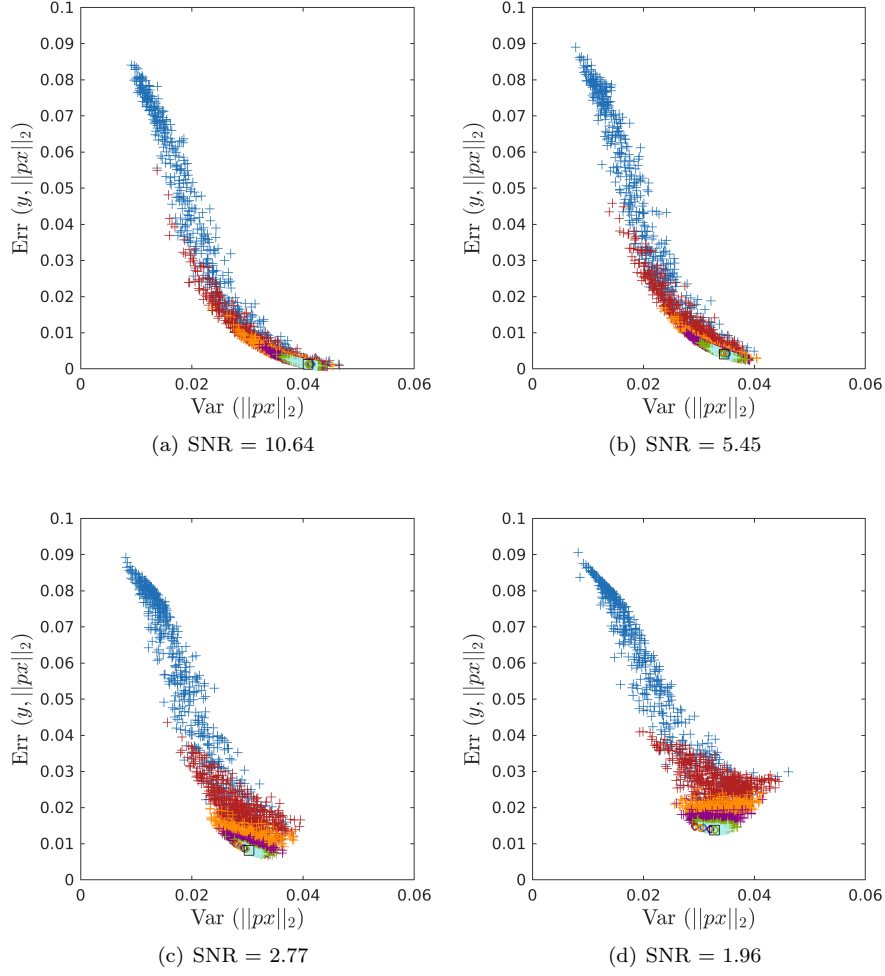


Figure 1: Scatter plots for the simulated data set with different noise levels, showing the L_2 reconstruction error (4) and variance (5) of the final images obtained by combining the coils with rSOS and compression by random projections and PCA in $\mathcal{G}_{k,32}$. A varying amount of Gaussian noise was added in each coil assuming no correlation [1]. The SNR value corresponds to the mean over all voxels in the reconstructed rSOS image volume without the linear compression.

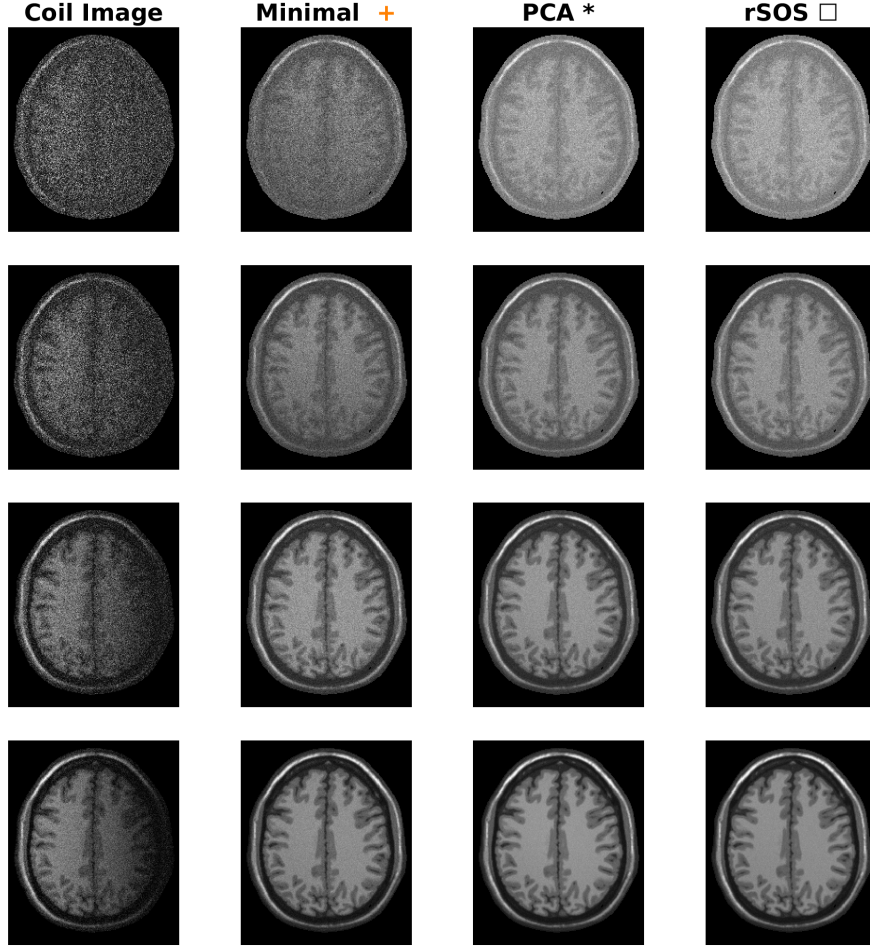


Figure 2: Cross-sectional image corresponding to the scatter plots in Figure 1 for the simulated data set with different noise levels. The left column shows the first channel of the simulated noisy 32-dim coil array before the coil combination step. The second column shows the final image obtained by rSOS including compression with the random projection $p \in \mathcal{G}_{28,32}$ that yields the minimal L_2 reconstruction error (4). The third column shows the image provided by rSOS without compression. The second last columns show the final images when using PCA in $\mathcal{G}_{1,32}$ and $\mathcal{G}_{4,32}$ for compression, yielding the highest mean SNR. We can directly see that the lowest L_2 error does not directly correspond to the highest SNR. Moreover, the visual evaluation suggests that the SNR describes the image quality more accurately.

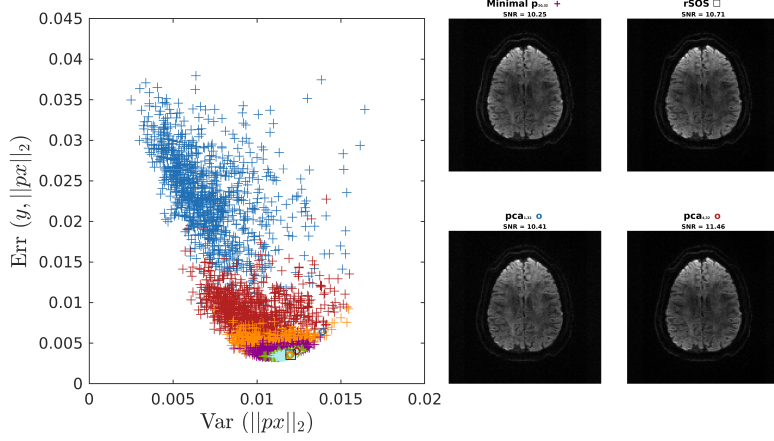


Figure 3: Left - Scatter plot for the in-vivo data set showing the reconstruction error (4) and variance (5) obtained by combining the coils with rSOS and compressing with random projections and PCA in $\mathcal{G}_{k,32}$. Right - Final cross-sectional images corresponding to the compression methods in the scatter plot. The SNR value corresponds to the mean over all voxels in the reconstructed image volume, see (6).

5 Discussion

In Figure 1 and 3 we illustrate the scatter plots of the described linear compression methods regarding the two data sets. To simplify the visualization we display the spaces $\mathcal{G}_{k,32}$ only for $k \in \{1, 4, 12, 20, 28, 31\}$. In all plots we can see that the smaller the dimension k , the more the projections are spread regarding $\text{Err}(y, \|px\|_2)$ and $\text{Var}(\|px\|_2)$; the reconstructions including random projections from $G_{1,32}$ are widely distributed, whereas using the projections in $G_{31,32}$ are clustered closely around the rSOS. The smaller the k , the more original information can be randomly dismissed: the preservation and loss of the original information varies less for image space data compressed by a random projection $p \in G_{31,32}$ in comparison to compression by $p \in G_{1,32}$.

For lower noise levels in the simulated data, Figure 1 (a)-(b), we can directly see that the correlation between the L_2 reconstruction error and the variance within the final combined images is very strong. Since variance can be interpreted as contrast in images with high SNR, it shows that high contrast directly relates here to a good reconstruction in the L_2 manner. However, the higher the noise level, the more noise is also described in the measured variance and therefore it cannot directly be interpreted as contrast any more. In the simulated data set, Figure 1 (c)-(d), this can be seen in the change of correlation behavior, where higher variance does not relate for all dimensions k to a lower reconstruction error. In Figure 1 (d) the projections from the spaces $G_{k,32}$ with $k = \{12, 20, 28, 31\}$ do not show any connection between the variance and L_2

error. This might happen because the high noise level in the original data influences the measure of variance strongly, merging the underlying meaning of contrast. Interestingly, for $k = 1$ there is still some negative correlation, indicating that the randomization leads to several poor reconstructions where noise is secondary in comparison to issues with contrast. The scatter plot corresponding to the in-vivo data (Figure 3) shows similar behavior. Because of this varying behavior, the variance itself does not act as a useful measure of reconstruction accuracy in this experimental setup.

We can see that in all experiments rSOS itself, as well as including PCA compression, yield low L_2 reconstruction errors, but are always outperformed by some random projections. Nevertheless, regarding SNR and visualization, these reconstructed images with random compression cannot compete with PCA compression or rSOS. That indicates some contradicting behavior, when measuring the reconstruction performance with the L_2 error versus the SNR. Indeed, when using PCA for the simulated data set, the L_2 error is the lowest in $G_{31,32}$, whereas highest SNR is always achieved by PCA in $G_{4,32}$, which contradictory yields the worst L_2 error. Also in the in-vivo data set the highest SNR was achieved by using PCA in $G_{4,32}$, which again does not correspond to the lowest L_2 error. Following the visual results it seems that the SNR describes here the visual performance better than the L_2 error. The compression with PCA in $G_{4,32}$ yields the highest SNR in all experiments and therefore outperforms rSOS without compression consistently. Compression with the standard PCA in $G_{1,32}$ yields predominantly better results than rSOS, but yields an insufficient visual result on the in-vivo data. Just using the first eigendirection yields a strong compression that loses too much important information here.

6 Summary

Based on random orthogonal projections we have shown a numerical investigation on reconstruction accuracy regarding rSOS coil combination with linear compression. Two different MR data sets were used for our experiments; a simulated T1 weighted data set with varying amount of noise and an in-vivo diffusion weighted data set. We used diverse measures of performance to evaluate the image quality of the reconstructions. For the lower noise levels in the simulated data, the L_2 reconstruction error yields strong correlation with the variance, but the behavior changes for higher noise levels and the in-vivo data. Moreover, measuring L_2 error and SNR acts contradictory in terms of optimality and in these cases we observe that the visual evaluation corresponds more to the SNR. The highest SNR values were achieved by incorporating PCA as compression before using rSOS, outperforming rSOS with no compression in all experiments. This clearly suggests to use PCA on image space data before computing the final coil combination with rSOS, yielding a beneficial denoising effect with higher SNR.

Future work shall include related experiments in the k-space before PI reconstruction, where linear compression is highly beneficial to save subsequent

computational costs.

7 Acknowledgment

The work was partly funded by the Austrian Marshall Plan Foundation and Vienna Science and Technology Fund (WWTF) through project VRG12-009.

References

- [1] S. Aja-Fernández and G. Vegas Sánchez-Ferrero. *Statistical Analysis of Noise in MRI*. 01 2016.
- [2] J. S. Brauchart, A. B. Reznikov, E. B. Saff, I. H. Sloan, Y. G. Wang, and R. S. Womersley. Random Point Sets on the Sphere - Hole Radii, Covering, and Separation. *Experimental Mathematics*, 27(1):62–81, 2018.
- [3] A. Breger, M. Ehler, and M. Gräf. Points on manifolds with asymptotically optimal covering radius. *Journal of Complexity*, 48:1–14, 2018.
- [4] A. Breger, J. Orlando, P. Harár, M. Doerfler, S. Klimescha, C. Grechenig, B. Gerendas, U. Schmidt-Erfurth, and M. Ehler. On orthogonal projections for dimension reduction and applications in augmented target loss functions for learning problems. *Journal of Mathematical Imaging and Vision (JMIV)*, 2019.
- [5] M. Buehrer, K. P. Pruessmann, P. Boesiger, and S. Kozerke. Array compression for MRI with large coil arrays. *Magnetic Resonance in Medicine*, 57(6):1131–1139, 2007.
- [6] Y. Chang and H. Wang. Kernel principal component analysis of coil compression in parallel imaging. *Computational and Mathematical Methods in Medicine*, 2018:1–9, 04 2018.
- [7] Y. Chikuse. *Statistics on special manifolds*. Lecture Notes in Statistics. Springer, New York, 2003.
- [8] A. Chu and D. Noll. Coil compression in simultaneous multislice functional MRI with concentric ring slice-GRAPPA and SENSE. *Magnetic resonance in medicine*, 76(4):1196–1209, oct 2016.
- [9] C. A. Cocosco, V. Kollokian, R. K.-S. Kwan, G. B. Pike, and A. C. Evans. Brainweb: Online interface to a 3d mri simulated brain database. In *NeuroImage*. Citeseer, 1997.
- [10] C. A. Cocosco, V. Kollokian, R.-S. Kwan, G. Pike, and A. Evans. Brainweb: Online interface to a 3d mri simulated brain database. *NeuroImage*, 5:425, 1997.
- [11] D. L. Collins, A. P. Zijdenbos, V. Kollokian, J. G. Sled, N. J. Kabani, C. J. Holmes, and A. C. Evans. Design and construction of a realistic digital brain phantom. *IEEE Transactions on Medical Imaging*, 17(3):463–468, June 1998.
- [12] S. Dasgupta and A. Gupta. An elementary proof of a theorem of johnson and lindenstrauss. *Random Structures & Algorithms*, 22(1):60–65, 2003.
- [13] B. Girod. Psychovisual aspects of image processing: What’s wrong with mean squared error? In *Proceedings of the Seventh Workshop on Multidimensional Signal Processing*, pages P.2–P.2, Sep. 1991.
- [14] M. Griswold. Basic Reconstruction Algorithms for Parallel Imaging. In *Parallel Imaging in Clinical MR Applications*, pages 19–36. Springer, Berlin, Heidelberg, 2007.
- [15] M. A. Griswold, P. M. Jakob, R. M. Heidemann, M. Nittka, V. Jellus, J. Wang, B. Kiefer, and A. Haase. Generalized autocalibrating partially parallel acquisitions (GRAPPA). *Magnetic Resonance in Medicine*, 47(6):1202–1210, 2002.
- [16] W. S. Hoge and J. R. Polimeni. Dual-polarity GRAPPA for simultaneous reconstruction and ghost correction of EPI data. *Magnetic Resonance in Medicine*, 76(1):32–44, 2016.

- [17] R. K. Kwan, A. C. Evans, and G. B. Pike. Mri simulation-based evaluation of image-processing and classification methods. *IEEE Transactions on Medical Imaging*, 18(11):1085–1097, Nov 1999.
- [18] G. Lemberskiy, S. Baete, J. Veraart, T. M. Shepherd, E. Fieremans, and S. Novikov. Achieving sub-mm clinical diffusion mri resolution by removing noise during reconstruction using random matrix theory. *Proceedings of the International Society for Magnetic Resonance in Medicine*, 27, 2019.
- [19] M. Lustig and J. M. Pauly. SPIRiT: Iterative self-consistent parallel imaging reconstruction from arbitrary k-space. *Magnetic resonance in medicine*, 64(2):457–471, aug 2010.
- [20] K. P. Pruessmann, M. Weiger, M. B. Scheidegger, and P. Boesiger. SENSE: Sensitivity encoding for fast MRI. *Magnetic Resonance in Medicine*, 42(5):952–962, 1999.
- [21] A. Reznikov and E. B. Saff. The Covering Radius of Randomly Distributed Points on a Manifold. *International Mathematics Research Notices*, 2016(19):6065–6094, 12 2015.
- [22] P. B. Roemer, W. A. Edelstein, C. E. Hayes, S. P. Souza, and O. M. Mueller. The NMR phased array. *Magnetic Resonance in Medicine*, 16(2):192–225, 1990.
- [23] N. J. Tustison, Y. Yang, and M. Salerno. Advanced normalization tools for cardiac motion correction. In O. Camara, T. Mansi, M. Pop, K. Rhode, M. Sermesant, and A. Young, editors, *Statistical Atlases and Computational Models of the Heart - Imaging and Modelling Challenges*, pages 3–12, Cham, 2015. Springer International Publishing.
- [24] J. Veraart, D. S. Novikov, D. Christiaens, B. Ades-Aron, J. Sijbers, and E. Fieremans. Denoising of diffusion mri using random matrix theory. *NeuroImage*, 142:394–406, 11 2016.
- [25] D. O. Walsh, A. F. Gmitro, and M. W. Marcellin. Adaptive reconstruction of phased array MR imagery. *Magnetic Resonance in Medicine*, 43(5):682–690, 2000.
- [26] J. Wang, Z. Chen, Y. Wang, L. Yuan, and L. Xia. A feasibility study of geometric-decomposition coil compression in mri radial acquisitions. *Computational and Mathematical Methods in Medicine*, 2017:1–9, 01 2017.
- [27] Z. Wang, A. C. Bovik, and L. Lu. Why is image quality assessment so difficult? In *2002 IEEE International Conference on Acoustics, Speech, and Signal Processing*, volume 4, pages IV–3313–IV–3316, May 2002.
- [28] Z. Wang, A. C. Bovik, H. R. Sheikh, and E. P. Simoncelli. Image quality assessment: from error visibility to structural similarity. *IEEE Transactions on Image Processing*, 13(4):600–612, April 2004.
- [29] T. Zhang, J. M. Pauly, S. S. Vasanawala, and M. Lustig. Coil compression for accelerated imaging with Cartesian sampling. *Magnetic Resonance in Medicine*, 69(2):571–582, 2013.

A Appendix

A.1 Zusammenfassung

Im Zeitalter der Digitalisierung ist Dimensionsreduktion von großer Bedeutung um hochdimensionale Daten zu komprimieren. Eine niedrigdimensionale Datenrepräsentation ermöglicht effiziente Berechnungen, soll aber gleichzeitig essenzielle Informationen für nachfolgende Bearbeitung bewahren. In meiner Dissertation leite ich mathematische Ergebnisse zu orthogonalen Projektoren und Dimensionsreduktion her und verwende diese in medizinischen Bildverarbeitungsproblemen mit klinischen Daten. Die theoretischen Grundlagen ermöglichen neue Analysen und die Verbesserung von Lernalgorithmen mit neuronalen Netzen.

Orthogonale Projektoren sind eine effiziente lineare Methode zur Dimensionsreduktion. Für die Berechnung wird nur eine Matrixmultiplikation benötigt und dadurch wird effiziente Dimensionsreduktion auf großen und sehr hochdimensionalen Datenmengen möglich. Wir leiten neue theoretische Resultate für Punktfolgen her, die den Raum der orthogonalen Projektoren (die Grassmann-Mannigfaltigkeit) asymptotisch optimal abdecken. Durch die numerische Konstruktion können wir die theoretischen Resultate illustrieren und verwenden diese Punktfolgen auch in den weiteren Analysen und Experimenten.

Um die Eignung einer Dimensionsreduktionsmethode zu evaluieren, müssen verschiedene Aspekte des Informationserhalts beachtet werden. Der orthogonale Projektor gegeben durch die populäre Hauptkomponentenanalyse (PCA) maximiert die Totalvarianz innerhalb der projizierten Daten. Besonders Aufgabenstellungen mit verrauschten Daten können von PCA profitieren, da das Rauschen durch das Entfernen von den Eigenrichtungen reduziert werden kann, die zu den niedrigeren Hauptkomponenten korrespondieren. Hingegen erhalten zufällige Projektoren mit hoher Wahrscheinlichkeit alle paarweisen Distanzen innerhalb einer Datenmenge. Diese Eigenschaft ist von Bedeutung für Probleme in denen kleine Veränderungen innerhalb der projizierten Daten wichtig sind. Ob die eine oder andere Eigenschaft wichtig ist, kommt auf die jeweilige Anwendung und den Zustand der Daten an. Wir beweisen, dass normalerweise die gleichzeitige Erhaltung der Totalvarianz und der paarweisen Distanzen nicht möglich ist. Numerische Experimente illustrieren die theoretischen Ergebnisse.

Im angewandten Teil meiner Dissertation arbeite ich mit klinischen Bilddaten von optischer Kohärenztomographie (OCT) und Magnetresonanztomographie (MRT). Die Forschung fand in Kollaboration mit dem Vienna Reading Center (VRC), Fakultät für Ophthalmologie, Medizinische Universität Wien und dem Laboratory of Mathematics in Imaging (LMI), Brigham and Women's Hospital, Harvard Medical School, Boston (MA), statt. Automatisierte Bildsegmentierung ist besonders wichtig für Krankheitserkennung, da in der Regel der ärztliche Alltag nicht genug Zeit lässt um alle Bildaufnahmen von Patienten zu evaluieren. Wir entwickeln basierend auf manuell annotierten OCT Daten verschiedene Methoden zur Bildsegmentierung mit Lernalgorithmen, die Dimensionsreduktion und orthogonale Projektoren verwenden. Unsere Methoden könnten die Quantifizierung in pathologischen Daten verbessern und könnten in Zukunft als Richtlinien für Ärzte dienen. In einer Rekonstruktionsanalyse mit MRT Daten verwenden wir orthogonale Projektoren um multidimensionale Bilddaten zu komprimieren. Wir vergleichen die resultierenden Bildrekonstruktionen und folgern Verbesserung durch die zusätzliche Anwendung von PCA.

A.2 Curriculum Vitae

Curriculum Vitae of Anna Breger, BSc MSc

Office: Oskar-Morgenstern-Platz 1, Room 05.126,
1090 Wien, Austria

Phone (office): +43-1-4277-50429

E-mail: anna.breger@univie.ac.at

Website: homepage.univie.ac.at/anna.breger/



Academic positions

02/2019 - 05/2019 *Visiting Research Fellow*
Laboratory of Mathematics in Imaging (LMI), Brigham and
Women's Hospital, Harvard Medical School, Boston, MA

03/2016 – present *Research Assistant and PhD candidate*
University of Vienna, Department of Mathematics

Employed in the WWTF project 'Computational harmonic analysis of high-dimensional biomedical data' in collaboration with the Medical University of Vienna. Research on orthogonal projections used for dimension reduction as data preprocessing for subsequent learning algorithms with applications on medical image data.

University education and degrees

University of Vienna, Austria (Department of Mathematics)

09/2013 - 12/2015 Master of Science (MSc) - awarded with distinction
Master of Science Program for Applied Mathematics

Main focus on applied image and signal processing, numerical mathematics and different programming languages with the master's thesis 'On image segmentation and applications in clinical retinal analysis'.

09/2009 - 09/2013 Bachelor of Science (BSc) - awarded with distinction

Semesters abroad

01/2013 - 06/2013	University of Leeds (UK), Dept. of Mathematics
Summer, 2012	University of Texas at Austin, Dept. of Mathematics

Teaching

03/2018 – 08/2018	Lecturer, University of Vienna, Dept. of Mathematics Teaching basic algorithms in the programming language Python.
09/ 2013 - 09/2014	Tutor, University of Vienna, Dept. of Mathematics Teaching how to use the software LaTeX and Mathematica.

Related professional positions

07/2014	Internship at ProFactor GmbH, Vienna Development of a Matlab toolbox for creating various image features to classify errors.
---------	---

Awards

2019	Marshall Plan Scholarship Grant for the research project 'Image analysis in medical applications' Laboratory of Mathematics in Imaging, Harvard Medical School, Boston, MA
2018	‘Best PhD project’ awarded at Biostec 2018, International Conference
2015	Nominated for the high potential programme ‘Naturtalente’ of the University of Vienna Successful completion of expertise workshops for management, leadership and personality skills.
2014	Merit Scholarship for master studies, University of Vienna
2012	Merit Scholarship for bachelor studies, University of Vienna

Language skills

German (native), English (fluent), Spanish (basic)

Additional computer skills

Python, Matlab, Java, C++ ; LaTeX, Mathematica, Amira (image processing)

International presentations

- 10/2019 A. Breger, J.I. Orlando, H. Bogunovic, S. Riedl, B. S. Gerendas, M. Ehler, and U. Schmidt- Erfurth
An amplified-target loss approach for photoreceptor layer segmentation in pathological oct scans.
Talk and poster presentation at the Ophthalmic Medical Image Analysis (OMIA) workshop, Shenzhen, China
Invited talk at the Workshop on Computational Medical Imaging and Artificial Intelligence (WCMIA), Hangzhou, China
- 05/2019 A. Breger, G. Ramos-Llorden, G. Vegas Sanchez-Ferrero, C.-F. Westin
Investigation of MRI reconstruction accuracy with orthogonal projections,
Poster presentation at 7th Annual Radiology Research Symposium, Brigham and Women's Hospital, Boston, MA
- 04/2019 A. Breger et al.
On orthogonal projections for dimension reduction,
Talk at Applied Chest Imaging Laboratory, Brigham and Women's Hospital, Harvard Medical School, Boston, MA
- 03/2019 A. Breger et al.
On orthogonal projections for dimension reduction and applications in variational loss functions for learning problems
Talk at Laboratory of Mathematics in Imaging (LMI), Harvard Medical School, Boston, MA
- 02/2019 A. Breger, M. Ehler
Dimension reduction in learning tasks
Talk at 90th Annual Meeting of GAMM, Vienna, Austria
- 06/2018 A. Breger et al.
Learning and dimension reduction in medical image analysis
Poster presentation at SIAM Conference on Imaging Science, Bologna, Italy
- 03/2018 A. Breger, M. Ehler
Specific orthogonal projections for dimension reducing preprocessing in learning tasks,
Talk at 89th Annual Meeting of GAMM, Munich, Germany

- 01/2018 A. Breger, B.S. Gerendas, U. Schmidt-Erfurth, M.Ehler
Dimension Reduction and Automated Evaluation in Retinal OCT volumes,
 Talk and poster presentation at 11th international joint conference on
 biomedical engineering systems and technologies, Funchal, Madeira
- 09/2017 A. Breger, M. Ehler
*Mathematical image processing for automated evaluation in Retinal OCT
 volumes*,
 Poster presentation at 1st Interscience Symposium of the Vienna Doctoral
 Schools, Vienna, Austria
- 07/2017 A. Breger, M. Ehler
Sampling in Grassmannians,
 Poster presentation at SampTA, 12th international Conference on Sampling
 Theory and Applications, Tallinn, Estonia
- 03/2017 A. Breger, M. Ehler, H. Bogunovic, S.M. Waldstein, A. Philip, U.
 Schmidt-Erfurth, B.S. Gerendas
*A supervised learning pipeline for quantification of retinal fluid in medical
 OCT images*,
 Talk at 88th Annual Meeting of GAMM, Weimar, Germany
- 12/2016 A. Breger, H. Bogunovic, B.S. Gerendas, U. Schmidt-Erfurth, M. Ehler
*An efficient learning pipeline for quantification of retinal fluid in spectral-
 domain optical coherence tomography images*.
 Poster presentation at 2nd IMA Conference on the Mathematical
 Challenges of Big Data, London, UK

Publications

- 2019 A. Breger, J. I. Orlando, H. Bogunovic, S. Riedl, B.S. Gerendas, M. Ehler,
 and U. Schmidt- Erfurth
*An amplified-target loss approach for photoreceptor layer segmentation in
 pathological oct scans*, [arXiv:1908.00764](https://arxiv.org/abs/1908.00764)
 Ophthalmic Medical Image Analysis (OMIA), Lecture Notes in Computer
 Science, vol 11855, Springer
- P.Harar, R. Bammer, A. Breger, M. Dörfler and Z. Smekal
*Machines listening to music: the role of signal representations in learning
 from music*, [arxiv:1903.08950](https://arxiv.org/abs/1903.08950)
 11th International Congress on Ultra Modern Telecommunications and
 Control Systems (ICUMT), IEEE (in press)

A. Breger, J. I. Orlando, P. Harar, M. Doerfler, S. Klimscha, C. Grechenig, B. S. Gerendas, U. Schmidt-Erfurth, and M. Ehler
On orthogonal projections for dimension reduction and applications in augmented target loss functions for learning problems, [arxiv:1901.07598](https://arxiv.org/abs/1901.07598)
Journal of Mathematical Imaging and Vision (JMIV), Springer

2018

A. Breger, M. Ehler, and M. Gräf
Points on manifolds with asymptotically optimal covering radius,
[arxiv:1607.06899](https://arxiv.org/abs/1607.06899)
Journal of Complexity, vol 48, Elsevier

2017

A. Breger, M. Ehler, M. Gräf, and T. Peter
Cubatures on Grassmannians: moments, dimension reduction, and related topics, [arXiv:1705.02978](https://arxiv.org/abs/1705.02978)
Compressed Sensing and its Application (Applied and Numerical Harmonic Analysis), Springer

A. Breger, M. Ehler, H. Bogunovic, S.M. Waldstein, A. Philip, U. Schmidt-Erfurth, and B.S. Gerendas
Supervised learning and dimension reduction techniques for quantification of retinal fluid in optical coherence tomography images,
[nature.com/articles/eye201761](https://www.nature.com/articles/eye201761)
Eye, Springer Nature

A. Breger, M. Ehler, and M. Gräf
Quasi Monte Carlo integration and kernel-based function approximation on Grassmannians, [arXiv:1605.09165](https://arxiv.org/abs/1605.09165)
Frames and Other Bases in Abstract and Function Spaces, Applied and Numerical Harmonic Analysis series (ANHA), Birkhauser/Springer

Preprints

2019

A. Breger, G. Ramos Llorden, G. Vegas Sanchez – Ferrero, W. S. Hoge, M. Ehler and C-F. Westin
On the reconstruction accuracy of multi-coil MRI with orthogonal projections, [arXiv:1910.13422](https://arxiv.org/abs/1910.13422)