

WENCESLAS OR NOT - FACE COMPARISON ON FRANA'S WENCESLAS DEPICTION IN THE WENCESLAS BIBLE

Heinz Hofbauer¹ • Julia Hintersteiner¹ • Manfred Kern¹ • Andreas Uhl¹

¹Paris Lodron University of Salzburg, Austria (`{hofbauer, uhl}@cs.sbg.ac.at` `{julia.hintersteiner, manfred.kern}@plus.ac.at`)

July 9, 2025

Abstract

Our goal is to use computer vision to find commonalities in how faces are painted in late medieval illustrated manuscripts on the example of the Wenceslas Bible. Specifically, for the master illustrator referred to as Frana we want to find out how many illustrators from his workshop contributed to the Bible. We will limit ourselves to depictions of Wenceslas and use face recognition methods to attempt to find stylistic differences which can identify individual illustrators. Further, there are a number of images that might be depictions of Wenceslas, but the actual association is tenuous. We will attempt to use the same face recognition methods to try and answer the question whether or not the images in dispute are actually of Wenceslas.

Contents

1	Introduction	2
2	Related Work	2
3	Experiments	3
3.1	Data and Groundtruth	3
3.2	Evaluating Algorithms	3
3.3	On <i>Wenceslas</i> ?	5
3.4	On the Number of Illustrators	5
4	Conclusion	6

1 Introduction

Our goal is to use computer vision, taking hints from face recognition, similar to what [1] did for renaissance painters, to find commonalities in how faces are painted in the Wenceslas Bible. The Wenceslas Bible a late 14th century (about 1390) translation of the Latin Vulgate, and is lavishly illustrated. The Wenceslas Bible is a valued cultural heritage item and is stored in the Austrian National Library, it consists of six manuscripts with shelf marks codex 2759 to 2764. Of art historic interest, and to better understand the effort involved in creating a work like the Wenceslas Bible, is to find out how many people were involved in illustrating the Bible, as well as which parts each person did. We differentiate between illustrators and master illustrators. The masters are differentiated by style and given names of convenience, because actual names are not recorded, typically based on a famous piece of art they did. Further, these masters are assumed to preside a workshop, so other members of the workshop which study under this master have potentially also contributed to the work as illustrators. In previous work there is an attribution of illustrations to master illustrators [2], however, this is not an assured connection since there are no historical records about the illustrators. One of the masters, Frana, signed at least some of his work so we will, for now, limit ourselves to work by this master, or rather masters workshop. We will use face recognition methods to assess the images and try to discern the painted images to form an opinion on the number of illustrators actually working on the Bible from the Frana-workshop. The basic assumption is that recurring characters, such as Wenceslas and the “bathmaid”, are painted similarly by the same illustrator. The illustrators likely never saw the king, and figures like the bathmaid are entirely fictional, we assume that an idealized version is painted, resulting in a similar outcome for the same illustrator. The assumption is that a similarity in face recognition thus indicates that the same illustrator created the assessed faces.

In addition to the number of illustrators from the Fran-workshop working on the bible there is another question when it comes to depictions of Wenceslas. There are a number of images of which it is uncertain, tagged as *Wenceslas?* in [2] whether they actually depict Wenceslas, see Figure 1 for an example.

We want to answer two questions. 1) Are the *Wenceslas?* persons actually Wenceslas or not? and 2) How many illustrators were involved in painting the Frana Wenceslas images?

2 Related Work

No work was done on face recognition in 14th century book illustrations but some on the more general topic of faces in art. Srinivasan et al. [1] used facial landmarks and paint styles to identify whether different Renaissance portraits are from the same painter. Similarly, Zhong [3] applied face biometric recognition to Song dynasty paintings as a further argument in art historical discussions, such as whether a painter had self inserted their likeness into a painting. Liang [4] performed age and gender classification of faces in Japanese art starting from

cropped faces, their result was that an ensemble of different CNNs worked best. Wechsler et al. [5] introduced a “faces in art” database with modern art, a different topic than ours, and they highlighted that different art styles impact face detection differently. While primarily about face detection a part of their conclusion is also that face authentication is not a solved problem.

Our primary assumption is that face recognition comparators can deal with faces in art. However, we have seen from the above works that this can be a difficult process and that what works for one art style, e.g. Modern Art, might not work in another, e.g. Renaissance portraits. Since we lack the amount of data (and the necessary ground truth) to train new or re-train existing methods we can only test a large number of face recognition tools from literature to find a method or set of methods which work well. To further augment the face recognition methods we will also use some texture related methods to see if they can help to augment face recognition methods.

Most face recognition implementations are from Deepface/Lightface¹, for a performance comparison on regular faces see [6]. The following face recognition methods are used. DeepID [7] creates a set of highly compact and discriminative features summarizing from multi-scale mid-level features of a hierarchy of deep ConvNets. VGG-Face[8] learns a face classifier with a softmax output layer, then removes this layer to produce a face embedding, and further fine-tunes the network with triplet loss to ensure that embeddings of the same identity are close and those of different identities are far apart in Euclidean space, optimizing performance for face verification and identification tasks. FaceNet[9] is a deep convolutional neural network which maps faces into a compact 128-dimensional Euclidean space, using a triplet loss function, where the distance between embeddings corresponds to face similarity. DLib[10] which uses a CNN to produce a 128-dimensional face descriptor with the Euclidean distance between a pair of descriptors as the biometric comparator. OpenFace[11] also uses the embedding of facial features, this time into a 128-dimensional hypersphere, allowing for the closeness-is-similarity calculations with spherical distances instead of Euclidean distances. ArcFace [12] uses an additive angular margin loss to enhance the discriminative power of feature embeddings, significantly improving both intra-class compactness and inter-class discrepancy for robust biometric face comparison. SFace[13] combines anchor-based and anchor-free methods within a unified network architecture and a sigmoid-constrained hypersphere loss function to prevent over fitting on pristine images, thereby making it robust to noisy real world images. GhostFaceNets[14] are lightweight face recognition models which generate additional feature maps efficiently through inexpensive linear transformations to achieve high accuracy with much lower computational cost than traditional CNNs. They use ArcFaces loss function.

The following methods are not based on faces but are based on color, image structure and visual similarity. The learned perceptual image patch similarity (LPIPS)[15] uses features from image classification for the calculation of a visual sim-

¹<https://github.com/serengil/deepface>



Figure 1: One of the questions we want to answer: Is this a depiction of Wenceslas (**left**). For comparison a Wenceslas image from the same page (**1st on right**) with two other Wenceslas images. Algorithmically the answer is ‘No’.

ilarity score between two images. Color coherence vectors (ccv)[16] are an improved color comparison tool, taking not only the color histogram of an image into consideration but also the spatial distribution and intermixing of colors. Local binary patterns (LBP)[17] calculate a histogram of structural elements in an image and a comparison calculates a similarity of structural elements between images. Hu-moments (HU)[18] are a set of seven numerical values derived from an image’s central moments that are invariant to translation, scale, and rotation, making them robust for comparing shapes regardless of their position, size, or orientation in an image.

3 Experiments

3.1 Data and Groundtruth

For the experiments we will only look at the first book, signature codex 2759, of the Wenceslas Bible as it is the most complete. Furthermore, to limit the unknowns we will only look at Frana since his works are best attributable as he signed some of his work, or rather signatures given to him to work on. The name Frana is from the signature but is commonly associated with František[2], one of Wenceslas court painters. The faces were hand annotated and tagged with information from the descriptions in [2], leading to 15 pictures of Wenceslas (*Wenceslas*). Further, there are two pictures which are tagged as *Wenceslas?* because they look like him but the association is uncertain. In our evaluation we added a third *Wenceslas?*, the illustration of a man in a bathing scene in the marginalia on folio 10v, i.e., the 10th sheet (folio) back side (verso shortened to ‘v’, front side would be recto shortened to r). The typical bathing scene in the marginalia depicts Wenceslas and is thus strongly associated with him, which is the reason we include this as a maybe.

3.2 Evaluating Algorithms

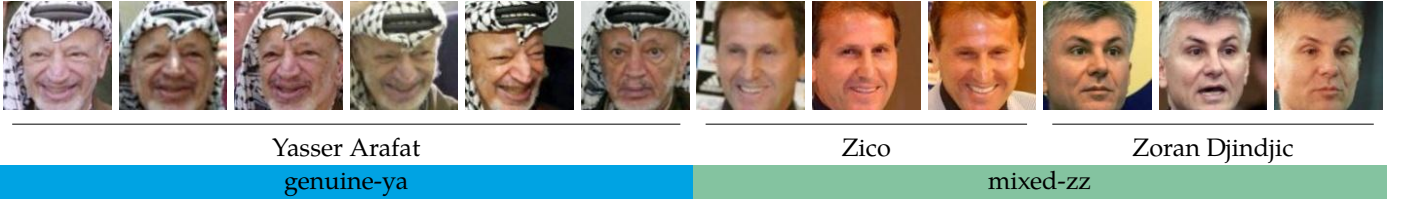
We have a number of algorithm which we can use. From prior work and literature we know that not all methods work equally well with any given art style. An evaluation is obviously called for, but we have a lack of ground truth on which to test the algorithms. What we want to know is if an algorithm can distinguish between imposter comparisons, i.e., faces do not belong to the same person, and genuine comparisons, i.e., the face belongs to the same person. However, the assumption is

that Wenceslas depictions painted by different illustrators can be differentiated by the face recognition, meaning that Wenceslas images from different illustrators should be counted as imposters. Generating real imposter comparisons is obviously not a problem, but since the number of illustrators for the Wenceslas images is in dispute we don’t know which comparisons are genuine and which would be counted as imposter.

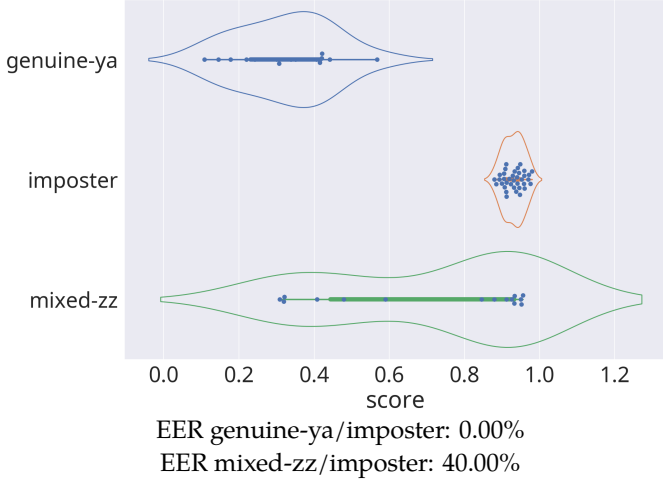
However, we can assume that some comparisons are genuine, i.e., at least in some instances an illustrator did more than one Wenceslas painting. Which would mean the resulting distribution would contain genuine and imposter comparisons and that should be detectable. The expected outcome can be illustrated on a created real world example, taking faces from the ‘Faces in the Wild’[19] dataset we construct two sets: (YA) containing 6 faces of a single person (Yassar Arafat) and one set (ZZ) containing images (3 each) of two unrelated persons (Zico and Zoran Djindjic). The first is clearly a full genuine set (what we would see if a single illustrator drew all of the Wenceslas images), the ZZ set is a mixed imposter/genuine set (what we would see if multiple illustrators, two in this case, drew some of the images), and a comparison between ZZ and YA is a set of imposter comparisons. Figure 2 shows setup (2a) and outcome with a face recognition method (2b), which more clearly separates the distributions, as well a texture classification based approach (2c) which shows a less ideal version with more overlap between distributions. The equal error rate (EER), i.e. the separation between the distributions at which the false positive and false negative rate are equal, is also given.

Translating to the Wenceslas Bible, our genuine set (or mixed set, depending on the number of illustrators) are the *Wenezsl* images of Frana. To form an imposter set we compare the *Wenceslas* images with another distinct image set, that of the *Bathmaid* another frequently occurring motive. The intra-*Bathmaid* comparisons are of no interest and thus skipped.

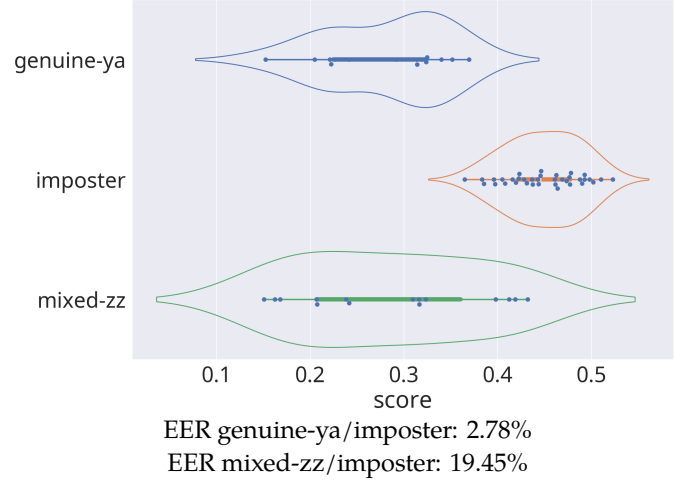
Two more settings should be taken into account, illustrated in Fig. 3. The first is the context of the face. If we cut out too close to the face the face comparison algorithms might have difficulty detecting the face so a bit of context would be beneficial. On the other hand, the texture around the face given by this context would be included in comparisons by the texture based metrics and have a detrimental influence. The second is orientation of the face, as the faces are often oriented up or down due to contents of the scene (looking to the heaven, lying down, dying and so on). Depending on the face recognition method they may have only a limited capacity to compensate



(a) Real world example of mixed versus genuine face comparisons.



(b) VGG-Face biometric comparison of faces.



(c) LPIS image classification feature based comparison of faces

Figure 2: Example of genuine, imposter and mixed distributions to show the likely outcome of the algorithm evaluation. Distributions as well as the resulting equal error rate (EER) between the genuine or mixed and imposter distribution is also given.

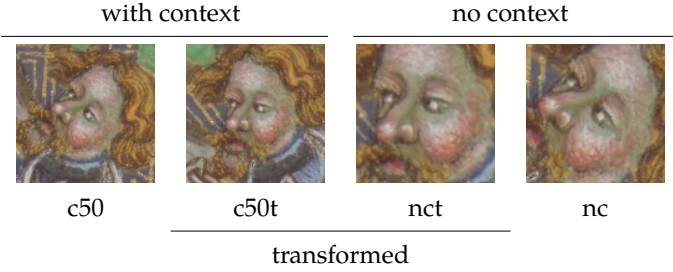


Figure 3: Example of test cases with a 50 pixel (*c50*) or without (*nc*) context around the face as well as unmodified or transformed to be upright (*t*).

for the orientation. As such we will test each algorithm on the regular image as well as a transformed (*t*) image where they faces are rotated upright. Likewise we will test on faces with no context (*nc*) or with a 50 pixel context on each side (*c50*).

Due to space limitations we cannot reproduce the large number of figures required to plot all outcomes, thus we will reduce each result to the equal error rate (EER). An EER of roughly 50% means that genuine and imposter distributions overlap completely. The farther removed from 50% the EER is the better the two distributions can be separated.

The results of the evaluation are given in Table 1 with relatively useless results ($\text{EER } 50\% \pm 10\%$) in red and the better results ($\text{EER} < 25\%$) in bold. The first point to make here is that some methods work, which in turn means the *Wenceslas* and *Bathmaid* images can be separated. This in turn means that

Table 1: The results of the evaluation of the algorithms on Frana’s *Wenceslas* (genuine comparisons) and *Wenceslas/Bathmaid* (imposter comparisons) distributions condensed to EER (in percent) with results close to 50% in red and ‘good’ results, assuming a mixed genuine set, in bold.

Algorithm		Equal Error Rate (EER) [%]			
		c50	c50t	nct	nc
face comparison	ArcFace	18.82%	38.13%	33.27%	21.88%
	DeepID	52.81%	47.53%	45.39%	61.95%
	Dlib	36.30%	47.08%	47.21%	34.86%
	Facenet	28.17%	29.49%	39.23%	27.86%
	GhostFaceNet	24.56%	37.79%	35.90%	25.64%
	OpenFace	47.61%	48.72%	47.83%	39.63%
	SFace	30.77%	48.65%	43.59%	28.40%
	VGG-Face	20.58%	33.03%	28.60%	16.05%
texture comparison	ccv	46.15%	47.67%	43.59%	38.01%
	hu	50.11%	50.17%	49.25%	48.97%
	lbp	42.38%	48.21%	50.53%	48.95%
	lpips	40.42%	42.86%	47.34%	46.15%

methods not working can simply not handle, this specific, art style facial images. With that said it’s pretty clear that texture based comparison does not work well, and we will disregard these methods for the rest of the paper. Overall we can also see that the transformation, to rotationally align faces, was overall detrimental to the results, likely due to the added transforma-

tion distortions and the methods innate ability to handle some degree of rotation. Interestingly, the 50 pixel context around the face only improves about half the face recognition methods. The three best metrics are ArcFace, GhostFaceNet and VGG-Face and the only metrics which fail more or less completely are DeepID and OpenFace. For completeness sake, and for comparison to our assumption about mixed genuine distributions (Fig. 2) the distributions of the three best methods are given in Figure 4. The results are as expected, in all cases there is a clear differentiation of the Wenceslas distribution and the imposter distribution, typically it is more compressed and only partially overlaps as we would expect of a mixed distribution. That the mixed distribution is so strongly shifted towards the imposter distribution, compared to the real world examples, is likely due to the difference in presentation if regards to the training sets of the recognition methods.

3.3 On Wenceslas?

There are three depictions, tagged *Wenceslas?*, in the Wenceslas Bible which can be argued to be Wenceslas but for various reasons there is also justifiable doubt. It should be noted, that it is assumed that more than one illustrator from the Frana-workshop worked on the Bibel, of which more later, so stylistic difference might well be at play here. However, Wenceslas should be clearly identifiable as Wenceslas, no matter the style or illustrator, as such we will disregard those influences and purely go by whether or not a depiction can be identified as Wenceslas. Algorithmically, we could query an algorithm to return the closest label from all other images and if it is Wenceslas we would count the *Wenceslas?* image as a depiction of Wenceslas also. Given our error rates that result would also be quite error prone. Instead we will use the three best algorithms from our evaluation, ArcFace and GhostFaceNet with 50 pixel context and VGG-Face with a close face crop, and we will use them for a majority voting style consensus finding. However, again the rank one label could be error prone, thus we collect the top 10 closest images and will use a majority voting on those.

The voting results are given in Table 2 with results ordered by number of votes, but grouped by algorithm and *Wenceslas?* image. The contentious *Wenceslas?* from folio 10v is also given in Figure 1 with a couple of acknowledged *Wenceslas* images for comparison.

The results of *Wenceslas?* 10v and 53v are relatively clear when it comes to the question Wenceslas or not, no biometric face comparison has a single Wenceslas image in the top 10. For *Wenceslas?* 10v we can directly compare to an image of Wenceslas on the same page, see Fig. 1, and we can see clear differences, structure of the nose, facial hair and hair color. Although, when looking at Wenceslas images from other pages these differences do not all apply. It should be noted that we included this image into the *Wenceslas?*, Theisen[2] did not, the algorithm clearly agrees with Theisen here.

As for *Wenceslas?* 53v, described as ‘crowned man enthroned in the initial D’[2], it is interesting to see that the voting of the algorithms strongly correlates the depiction with god (*God*). An association not entirely unlikely when consid-

Table 2: Results of the voting based on the rank 10 for the three contentious Wenceslas depictions. Page reference, as folio number of cod.2759, as well as a reference image of the Wenceslas depiction are given for reference. The ten labels are given per algorithm.

	10v Wenceslas?		53v Wenceslas?		93r Wenceslas?	
	label	count	label	count	label	count
ArcFace	Moses	3	Bathmaid	3	Wenceslas	3
	Aaron	2	God	3	Bathmaid	2
	Bathmaid	1	Moses	2	Angel	2
	God	1	Josef	1	Aaron	1
	King Balak	1	Zippora	1	Samson	1
	Chimera	1			Wildman	1
	Samson	1				
GhostFaceNet	Aaron	3	God	4	Aaron	3
	Moses	3	Aaron	3	Bathmaid	3
	God	1	Bathmaid	1	Moses	3
	Jesus	1	Moses	1	King Balak	1
	Pharao	1	Pharao	1		
	Wildman	1				
VGG-Face	Moses	4	God	5	Moses	3
	God	3	Moses	2	Bathmaid	2
	Aaron	1	Josef	1	Aaron	1
	Eljasaf	1	Lot	1	Eljasaf	1
	King Balak	1	Pharao	1	God	1
					Wenceslas	1
					Wildman	1

ering that an illustrator from the clergy should draw a regal figure in the Bible.

The case for *Wenceslas?* 93r is less clear, with ArcFace clearly voting for a depiction of Wenceslas, although, equally clearly GhostFaceNet and VGG-Face vote against. Going by illustration there is also a stronger connection to Wenceslas as the tent the man is looking out of is decorated with the letters “w” and “e”, which are strongly associated with Wenceslas throughout the Bible. The letter “w” is commonly assumed to be a monogram of Wenceslas[2] and while the meaning of the letter “e” is less clear[20], it’s co-occurrence with a Wenceslas depiction in the codex is frequent.

3.4 On the Number of Illustrators

The number of illustrators associated with a given masters name, Frana in this case, is unknown, but the common conception is that the ‘named’ masters preside over a workshop. Their students would work on the same projects and often fill in sketches done by the master. Given that they study under the named master their style and technique will be very close to the master, as such a difference, as in who actually painted an image, is very hard to detect. We use agglomerative clustering[21] to display the relative closeness of images and how they would be combined into clusters, this can be nicely displayed as a dendrogram. The basic idea is that each image is it’s own cluster and each join combines two clusters together, the display is then the distances when the join occurs on the y-axis and the formed clusters on the x-axis. A good example is again our real world example from Section 3.2 which is given

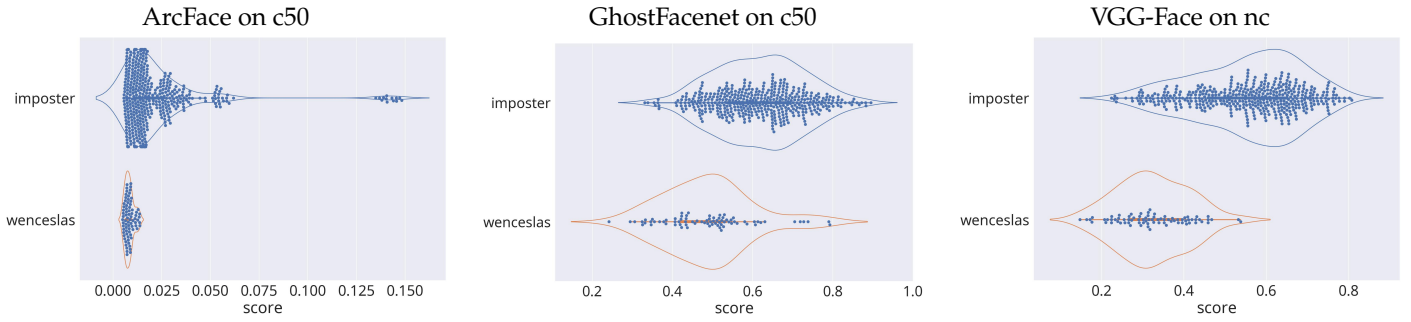


Figure 4: Genuine (inter *Wenceslas*) and imposter (*Wenceslas* vs. *Bathmaid*) distributions of the three best face recognition methods.

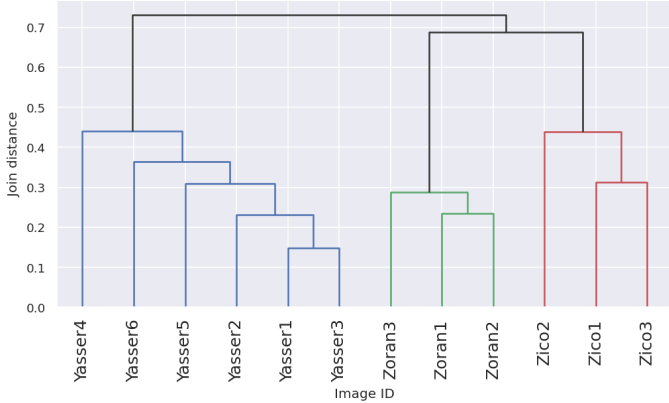


Figure 5: Example of a dendrogram with real world images (c.f. Section 2a). Distances at which a join occurs are given on the y-axis and image labels are given on the x-axis. The groupings into three clusters, corresponding to the three persons, can be clearly seen. Optimally the join in real clusters (colored here) and the joins between clusters (in black), are nicely separated by a large space on the y-axis.

in Figure 5. The clustering is done on the basis of the Euclidean distance in $\mathbb{R}^3$ where each dimension corresponds to the distance of one of the three algorithms (ArcFace, GhostFaceNet and VGG-Face).

The dendrogram for Frana’s *Wenceslas* images, with the actual images, are given in Figure 6. At first sight the three groupings, in color, make sense in that there is a difference between them. The largest group of images, in red, have a distinctly more coarse style, with brushstrokes being visible on the skin, as opposed to the other which were smoother. This alone does not guarantee that a different illustrator did them, as smoothness might have been sacrificed for speed. The difference between the other two large groups is less clear, but the blue (99r, 89r and 106r) group has an overall higher contrast in the facial area than the green group, with maybe the exception of 106r. The somewhat split 10v image is difficult to assess. Algorithmically it sits between the green and red groups but closer to the red, and it can be argued that it is rather similar to 51r from that group. However, overall the smoothness of 10v is very much like the paintings in the green group rather than the more coarse finish of the red group. The second non-grouped image is 62r, and the problem here is that the whole area of jaw and mouth seem to be drawn in a different orien-

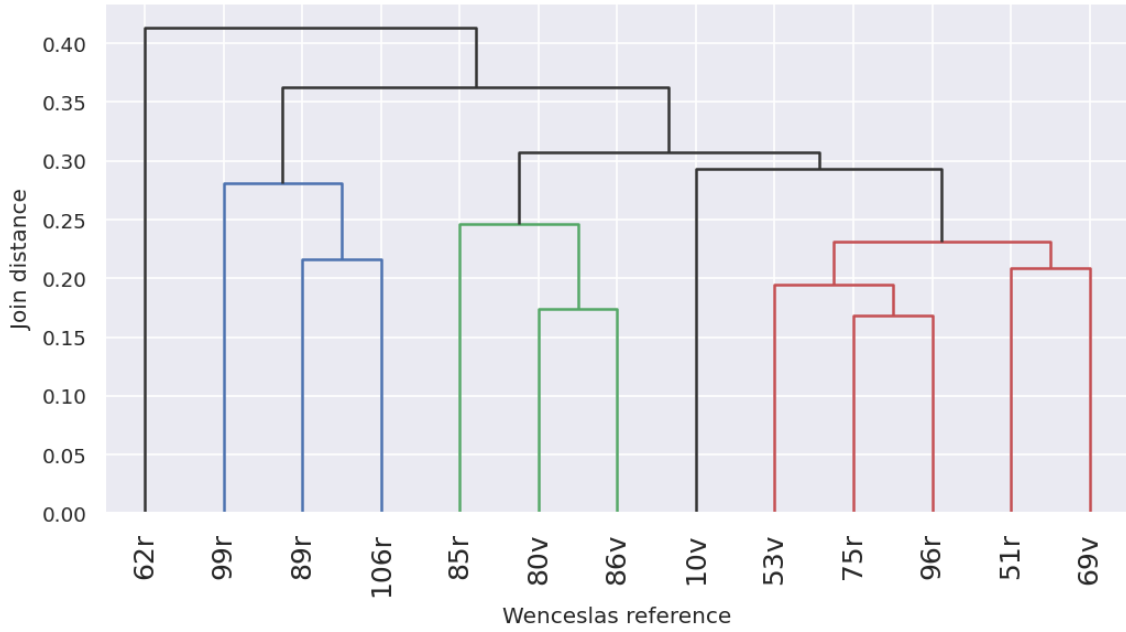
tation than the rest of the face. This likely throws off the face recognition algorithms. Subjectively, the placement of 62v is hard, it is very similar in style to 85r or 89r. Contrast wise it seems to have more bold colors than 85r but less value contrast than 89r. When looking at it however from the perspective of how many illustrators, we can state that it belongs to one or the other and is certainly not a different illustrator.

Are there at least two illustrators then? A likely conclusion, specifically because the difference between the blue and green groups are contrast, and typically novice painters lack the experience and confidence to utilize high contrast and the blue group has a markedly higher contrast in both value and hue than the green group. In the end this leaves the question are there three illustrators or two? An argument can be made for both, depending if the coarser drawing style of the red group is due to time constraints or personal style. Normally this a question that could not be answered, however, the algorithm used to separate this group is based on a face recognition, and while the algorithms might to some extent pick up on style as it influences the presentation of the face, the separation is quite clear. So from a biometric standpoint we have three groups of similar faces which would suggest that the painting style in group three is a personal one.

4 Conclusion

We set out to answer two questions with the help of face recognition methods. The first was regarding certain images are actual depiction of *Wenceslas*. Two, of which one was added by the authors due of the frequent association of *Wenceslas* with bathing scenes, were algorithmically rejected to be *Wenceslas*. The third, which also showed other small scene elements associated with *Wenceslas*, was voted 2:1 against by the algorithms, a decidedly less clear outcome.

The second question was about the number of illustrators from the Frana-workshop which worked on the illustrations, based on the *Wenceslas* images. From a face biometric standpoint, the images can be grouped into roughly three groups, with a very distinct outlier, which also correspond to one quite distinct and one smaller stylistic difference. So our tentative assessment would be that three illustrators were working on the *Wenceslas* images produced by the Frana-workshop.



(a) Dendrogram of join decisions for Frana's Wenceslas images.



(b) Wenceslas images with page reference, ordering and color matching the dendrogram.

Figure 6: Dendrogram of the agglomerative clustering of Frana's Wenceslas images and matching images. The clustering is done on the basis of the Euclidean distance in $\mathbb{R}^3$ where each dimension corresponds to the distance of one algorithm.

Acknowledgment

This project received funding from the Salzburg State Digital Humanities project “Digitalisation in the Humanities, Social and Cultural Sciences (HSC)” [22].

References

- [1] R. Srinivasan, C. Rudolph, and A. K. Roy-Chowdhury, “Computerized face recognition in renaissance portrait art: A quantitative measure for identifying uncertain subjects in ancient portraits,” *IEEE Signal Processing Magazine*, vol. 32, no. 4, 2015. doi: [10.1109/MSP.2015.2410783](#) (cit. on p. 2).
- [2] M. Theisen and U. Jenni, *Mitteleuropäische Schulen IV (ca. 1380–1400); Textband: Hofwerkstätten König Wenzels IV. und deren Umkreis*. 2014. doi: [10.26530/oapen_507994](#) (cit. on pp. 2, 3, 5).
- [3] G. Zhong, “A computer vision-aided analysis of facial similarities in song dynasty imperial portraits,” *Electronic Imaging*, vol. 35, no. 13, 2023. doi: [10.2352/EI.2023.35.13.CVAA-212](#) (cit. on p. 2).
- [4] Z. Liang, “Face recognition from art face images based on deep learning,” in *Proceedings of the 4th International Conference on Big Data Research*, ser. ICBDR ’20, 2021, isbn: 9781450387750. doi: [10.1145/3445945.3445963](#) (cit. on p. 2).
- [5] H. Wechsler and A. S. Toor, “Modern art challenges face detection,” *Pattern Recognition Letters*, vol. 126, 2019, issn: 0167-8655. doi: [10.1016/j.patrec.2018.02.014](#) (cit. on p. 2).
- [6] S. Serengil and A. Ozpinar, “A benchmark of facial recognition pipelines and co-usability performances of modules,” *Journal of Information Technologies*, vol. 17, no. 2, 2024. doi: [10.17671/gazibtd.1399077](#) (cit. on p. 2).
- [7] Y. Sun, X. Wang, and X. Tang, “Deep learning face representation from predicting 10,000 classes,” in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2014 (cit. on p. 2).
- [8] O. M. Parkhi, A. Vedaldi, and A. Zisserman, “Deep face recognition,” 2015 (cit. on p. 2).
- [9] F. Schroff, D. Kalenichenko, and J. Philbin, “Facenet: A unified embedding for face recognition and clustering,” in *2015 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2015. doi: [10.1109/CVPR.2015.7298682](#) (cit. on p. 2).
- [10] D. E. King, “Dlib-ml: A machine learning toolkit,” *Journal of Machine Learning Research*, vol. 10, 2009 (cit. on p. 2).
- [11] B. Amos, B. Ludwiczuk, and M. Satyanarayanan, “Openface: A general-purpose face recognition library with mobile applications,” CMU-CS-16-118, CMU School of Computer Science, Tech. Rep., 2016 (cit. on p. 2).
- [12] J. Deng, J. Guo, N. Xue, and S. Zafeiriou, “Arcface: Additive angular margin loss for deep face recognition,” in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2019 (cit. on p. 2).
- [13] Y. Zhong, W. Deng, J. Hu, D. Zhao, X. Li, and D. Wen, “Sface: Sigmoid-constrained hypersphere loss for robust face recognition,” *IEEE Transactions on Image Processing*, vol. 30, 2021. doi: [10.1109/TIP.2020.3048632](#) (cit. on p. 2).
- [14] M. Alansari, O. A. Hay, S. Javed, A. Shoufan, Y. Zweiri, and N. Werghe, “Ghostfacenets: Lightweight face recognition model from cheap operations,” *IEEE Access*, vol. 11, 2023. doi: [10.1109/ACCESS.2023.3266068](#) (cit. on p. 2).
- [15] R. Zhang, P. Isola, A. A. Efros, E. Shechtman, and O. Wang, “The unreasonable effectiveness of deep features as a perceptual metric,” in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2018 (cit. on p. 2).
- [16] P. Hiremath and J. Pujari, “Content based image retrieval using color, texture and shape features,” in *15th International conference on advanced computing and communications (ADCOM 2007)*, IEEE, 2007 (cit. on p. 3).
- [17] T. Ahonen, A. Hadid, and M. Pietikainen, “Face description with local binary patterns: Application to face recognition,” *IEEE Transactions on Pattern Analysis & Machine Intelligence*, no. 12, 2006 (cit. on p. 3).
- [18] Z. Huang and J. Leng, “Analysis of hu’s moment invariants on image scaling and rotation,” in *2010 2nd International Conference on Computer Engineering and Technology*, vol. 7, 2010. doi: [10.1109/ICCET.2010.5485542](#) (cit. on p. 3).
- [19] G. B. Huang, M. Ramesh, T. Berg, and E. Learned-Miller, “Labeled faces in the wild: A database for studying face recognition in unconstrained environments,” University of Massachusetts, Amherst, Tech. Rep. 07-49, 2007 (cit. on p. 3).
- [20] D. Gresch, “Das »e« in der wenzelsbibel,” *Kunstchronik. Monatsschrift für Kunstwissenschaft*, vol. 57, no. 3, 2004. doi: [10.11588/kc.2004.3.81507](#) (cit. on p. 5).
- [21] D. Müllner, *Modern hierarchical, agglomerative clustering algorithms*, 2011. arXiv: [1109.2378 \[stat.ML\]](#) (cit. on p. 5).
- [22] E. K. des Fachbereichs Germanistik der Universität Salzburg und der Österreichischen Nationalbibliothek. “Die Wenzelsbibel – Digitale Edition und Analyse.” version 2.2.0. Webpage: <https://edition.onb.ac.at/wenzelsbibel>. (2024), (visited on 01/25/2024) (cit. on p. 8).