

Blind Text Image Super-resolution for Enhancing the Readability of Fragments Hidden in Book Bindings

Baharan Pourahmadi¹ , Charlotte Eppe¹ , and Mads Toudal Frandsen² 

¹ Department of Culture and Language, University of Southern Denmark, Odense, Denmark

² Department of Physics, Chemistry and Pharmacy, University of Southern Denmark, Odense, Denmark

Abstract

Fragments reused in bookbindings are a crucial source for studying pre-modern book culture, representing books that were no longer read, and thus repurposed for their material value. The texts on these fragments are often hidden beneath pastedowns, scraped off, or otherwise obscured. Recovering their content is essential for paleographic and book historical analysis. Imaging methods such as Hyperspectral Imaging (HSI), Multispectral Imaging (MSI), and Ultraviolet Reflected (UVR) photography, combined with post-processing, are commonly used to reveal hidden text. However, the resulting images are often blurry and of low resolution, limiting their readability by both Optical Character Recognition (OCR) systems and human experts. This paper explores the use of deep generative models for super-resolving Latin text in fragment images, with a focus on fragments from the Herlufsholm Collection at the University of Southern Denmark. This collection contains numerous books with fragments—a valuable but understudied resource for investigating Scandinavian book history. We evaluate the effectiveness of text-specific super-resolution models in enhancing legibility and demonstrate their potential to support fragmentological research by making unreadable text accessible for scholarly analysis. The relevant material is available at: https://github.com/bhrnprhmd/Text_image_super_res_folatin_CHR2025.

Keywords: Text image super-resolution, Historical document analysis, Book fragments, Deep generative models

1 Introduction

In the Middle Ages and the Early Modern period, cutting up old, no longer needed books and finding other uses for the material was common practice. Parchment and paper were valuable commodities, expensive and time consuming to produce, but also robust and versatile. Obsolete literature was often repurposed for making covers for other books. Such fragments are an important facet of medieval and early modern book culture, presenting a reservoir for potentially otherwise unknown texts, as well as illuminating the life cycles of books [22].

However, in many cases, the damage done by the process of reuse impedes the legibility, and thus the study of the texts they contain. Two such examples are Herlufsholm 250.8 and Herlufsholm 24.1. They take their shelfmarks from Herlufsholm Skole, a Danish elite school founded in 1565. The institution accumulated a substantial book collection over the centuries [8; 19]. In the 1960s, these books were sold to the University of Southern Denmark, where the 40,000 volumes form the core of the university library’s historical holdings. It contains many volumes bound in reused parchment and paper which have aroused scholarly interest in recent years [9; 10].

Herlufsholm 250.8 is an Early Modern medical text, the *Progymnasmata* by French physician Jacques Aubert, printed in Basel in 1579. It is bound in a contemporary binding that was colored

Baharan Pourahmadi, Charlotte Eppe, and Mads Toudal Frandsen. “Blind Text Image Super-resolution for Enhancing the Readability of Fragments Hidden in Book Bindings.” In: *Computational Humanities Research 2025*, ed. by Taylor Arnold, Margherita Fantoli, and Ruben Ros. Vol. 3. Anthology of Computers and the Humanities. 2025, 1369–1382. <https://doi.org/10.63744/Hgeypr5GUQfb>.

© 2025 by the authors. Licensed under Creative Commons Attribution 4.0 International (CC BY 4.0).

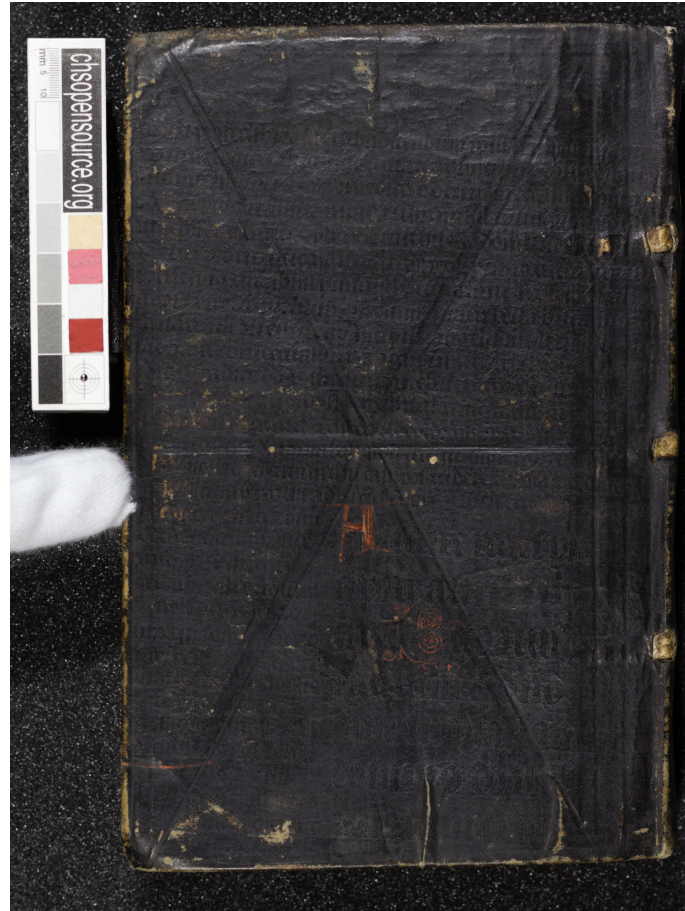


Figure 1: Back cover of Herlufsholm 250.8 under visible light, rotated 180 degrees to show fragment text correctly aligned.

black in order to obscure the fact that the parchment had already been used once. The varnish appears to be a thin, but opaque layer, obscuring the dark ink used for most textual elements well. Nonetheless, the text on the fragment has been identified as a passage from the Vulgate bible with glosses by Origen, a theologian of late antiquity, and dated to around 1300 [10].

Herlufsholm 24.1 consists of three printed volumes (Jacques Salian's *Annales Ecclesiastici*, printed in Cologne 1620–1624) bound in parchment over paste laminate boards made from paper waste with printed text on it. Herlufsholm 24.1 has recently been the subject of another case study [20], which states that the bindings were likely made in Denmark with local printer's waste.

In both cases the texts on the fragments are obscured and difficult to read. Developing a non-invasive method for improving legibility of such fragments can enable closer study in order to situate the fragments in the historical contexts in which they were produced and subsequently reused.

These books and fragments are historically valuable, and therefore cannot be physically altered or sampled for further analysis. Non-destructive imaging techniques such as Hyperspectral Imaging (HSI), Multispectral Imaging (MSI), and Ultraviolet Reflected (UVR) photography are widely used to recover obscured text without damaging the material [1; 3; 5; 23]. While these methods have shown success in many cases, severe degradation, such as dark varnishes, pigment bleed, or the presence of overlaid paper, often leads to light scattering and low-contrast images. As a result, the retrieved text is frequently of poor quality: unreadable by automated transcription models and, in extreme cases, even by human experts.



Figure 2: a) Herlufsholm 24.1 (RGB image captured with a standard camera). b) Hidden fragments are partially revealed at the edge of the binding.

Previous studies have explored the enhancement of degraded historical documents [7], but significantly less attention has been given to reused fragments, especially those reused in book bindings. The types of degradation targeted in earlier work are often less complex than those encountered in fragmentary materials, which may involve multiple layers, opacity from overlaid paper, or staining and abrasion from historical handling. These challenges call for more sophisticated restoration techniques. One of the main obstacles in restoring historical text images is the lack of representative datasets. In most cases, high-resolution ground truth images of the degraded texts are unavailable, making it difficult to train supervised models effectively. As a result, researchers often rely on synthetically generated datasets, pairing artificially degraded images with their clean counterparts. This approach requires a carefully designed degradation model that realistically captures the visual complexity of historical damage. Furthermore, because the goal of restoration in this context is to recover not just image quality but textual fidelity, any restoration method must ensure that character shapes, spacing, and stylistic features are preserved. These factors are essential for both human readability and for enabling downstream tasks such as Optical Character Recognition (OCR) or paleographic analysis.

Super-resolution techniques have previously been applied to text images and have shown promising results, particularly when using deep generative models such as MARCONet [16]. However, their applicability to historical fragments, especially those affected by physical layering or chemical darkening, remains largely unexplored. This study aims to investigate the potential of super-resolution approaches in the historical fragments context.

The following contributions are presented:

- 1) A blind super-resolution model designed for text image restoration was trained on Latin script data, and adapted to the characteristics of historical fragmentary texts.
- 2) To evaluate the effectiveness of the model in improving the readability of Latin text images, it was tested on synthetically degraded test samples, and the relevant image quality and readability metrics were measured. This quantitative evaluation was complemented by a qualitative assessment on a limited number of samples to visually examine the image quality of the super-resolved

outputs.

3) The trained model was applied to fragments from Herlufsholm 24.1 and 250.8, demonstrating measurable improvements in legibility and contributing to ongoing paleographic investigations of these materials.

2 Blind text super-resolution

Blind text image super-resolution (SR) is a specialized subdomain of image restoration focused on reconstructing high-resolution (HR) text images from low-resolution (LR) inputs that have been degraded by unknown processes. The term blind indicates that the degradation function is not known in advance. This is especially relevant for historical fragments, where degradation tends to be complex, spatially inconsistent, and often caused by a combination of factors such as pigment bleed-through, surface abrasion, layered materials (e.g., pastedown paper), and dark varnishes. Because the exact causes and extent of degradation in these documents cannot be easily modeled, blind SR approaches are more suitable than methods that assume a fixed or known degradation process.

Text image super-resolution differs fundamentally from general image SR, which primarily focuses on enhancing perceptual detail in natural scenes. In the context of text, the central goal is to improve readability and character clarity, which is essential for downstream applications such as OCR. Restoration methods must ensure high text fidelity and stylistic realism, as even minor distortions can render characters unrecognizable or ambiguous. Compared to natural image restoration, text image SR presents unique challenges. While minor imperfections may be acceptable in natural images, textual data demands precise structural recovery. Models must preserve fine character details and accurately reproduce typographic features to maintain the text's integrity. In response to these challenges, recent developments in text image restoration have introduced advanced architectures, including diffusion models and Transformer-based approaches, which have demonstrated significant improvements in handling complex degradations and preserving fine-grained textual features [2; 17]. Among these, MARCONet [16] has shown promising results, particularly in restoring low-resolution Chinese text images.

In this study, MARCONet was selected due to its demonstrated effectiveness and architectural suitability. This model enhances blind super-resolution (SR) of text images by leveraging structure priors learned from a StyleGAN [12; 13] trained on character images. A codebook constrains the generative space of the Generative Adversarial Network (GAN) [6], mapping each character to a specific latent code, while font style is controlled via the GAN's latent space. Once trained, the GAN provides intermediate features that serve as structure priors for characters. A Transformer-based encoder predicts the font style, character positions, and corresponding codebook indices from a LR text image. These priors are then used by a text SR network to guide the super-resolution of each character, aligning them via predicted bounding boxes.

To adapt the model for use on historical fragments, it was retrained on a dataset of Latin text, enabling its application to early modern and medieval script recovery tasks.

3 Training MARCONet on Latin text

3.1 Data preparation

To focus this study on Latin characters, a large corpus of Later Latin words was first compiled. A diverse range of historical Latin scripts was represented by using eight fonts, including Fraktur, Meyene, Schwabacher, and Times New Roman. To simulate realistic historical conditions, background textures were patched from images of actual historical documents and fragments. These backgrounds were then used to synthesize the final high-resolution (HR) text images, following a

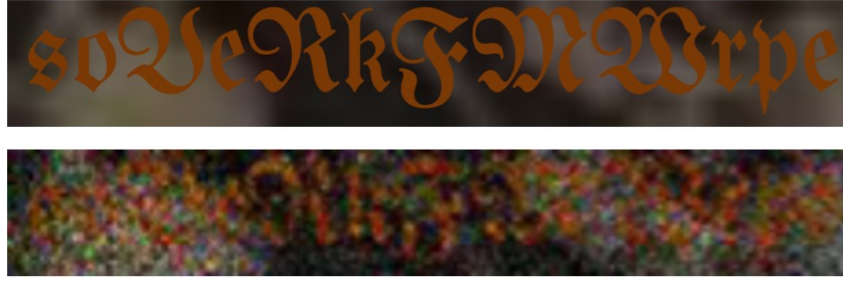


Figure 3: Example of Latin HR-LR synthetic data sample pair.

process similar to that of MARCONet. To generate the corresponding low-resolution (LR) samples, degradation pipelines from BSRGAN [25] and Real-ESRGAN [24] were applied to the HR images. The degradation pipeline consists of the following:

- 1) Random blurring: A Gaussian blur with a randomly selected kernel size and standard deviation is applied.
- 2) Downsampling: The blurred image is downsampled using a random scaling factor (typically between 1/2 to 1/4) and interpolation methods (bicubic, bilinear, or nearest neighbour), simulating spatial resolution loss.
- 3) Noise addition: Gaussian and Poisson noise are added to mimic sensor and environmental noise.
- 4) Compression artifacts: JPEG compression with randomly selected quality settings introduces compression artifacts.
- 5) Final upsampling: The degraded image is upsampled back to the original resolution to form the LR image.

Figure 3 illustrates an example of an LR–HR image pair produced through this pipeline.

The synthetic dataset of Latin HR-LR pairs consist of 10,000 samples.

3.2 Training setup

The training procedure closely followed the specifications provided in the original MARCONet paper. The model pre-trained on Chinese text was used to initialize the network weights. Fine-tuning was then performed on the custom Latin dataset described earlier. A batch size of 2 was used during training. The Adam optimizer [15] was employed, with the learning rate set to 1×10^{-4} . All training was conducted using PyTorch, on two Tesla A100 GPUs.

4 Evaluation setup

4.1 Quantitative evaluation

To evaluate the model’s effectiveness in super-resolving Latin text images, we first tested it on synthetically generated samples. Because the original appearance of the text on real fragments is unknown, and since the fragments themselves cannot be physically separated, no ground truth exists for the real degraded samples.

A dataset consisting of 1,000 synthetic Latin text image samples was created, following the same procedure used for training data generation described in section 3.1. Each sample contains a pair of low-resolution and high-resolution images.

To quantitatively assess the performance of the trained model in super-resolving Latin text images, we measured standard image quality metrics, namely the Peak Signal-to-Noise Ratio (PSNR)

and the Structural Similarity Index Measure (SSIM), on super-resolved test images (relative to the high-resolution ground truths). To further quantify the readability of the reconstructed text images, we evaluated the OCR performance using the Character Error Rate (CER).

The Kraken [14] recognition model (i.e., a neural network-based OCR system designed for historical and non-standard scripts) was used to recognize the characters in both the super-resolved and low-resolution samples. The recognized characters from the high-resolution samples were used as the ground truth text.

As presented in Table 1, the CER for the super-resolved samples is lower than that of the low-resolution samples, indicating an improvement in character readability after applying the super-resolution model. The PSNR and SSIM values reported in the original MARCONet paper [16] (28.32 dB and 0.941, respectively) are slightly higher, likely due to the smaller size of our training dataset.

SSIM	PSNR (dB)	CER (SR) (%)	CER (LR)(%)
0.921	26.98	35.54	88.04

Table 1: Quantitative evaluation of the super-resolution model using image quality (PSNR, SSIM) and readability (CER) metrics on synthetic Latin text images.

4.2 Qualitative comparison

To simulate the effect of overlaid paper often encountered in historical bindings, we fabricated test samples by placing thin Washi (Japanese) paper over printed text. Washi paper is commonly used in conservation practices due to its transparency and structural properties. To reveal the underlying text in these samples, HSI was employed. Figure 4 shows how the addition of this paper layer significantly degrades the final image.

The trained model was then used to restore text that had been degraded at three different levels: with one layer of paper, two layers, and three layers placed on top of the printed text. This was done to visualize how well the model can handle increasing levels of visual obstruction. The qualitative comparison is presented in Figure 4. In addition to these real degradation scenarios, the model was also tested on an artificially degraded sample to further assess its performance (Figure 5).

The CER was measured on the fabricated LR inputs and SR outputs (samples from Figure 4). As shown in Table 2, the readability of super-resolved samples have increased compared to the low-resolution inputs.

	Output (LR)	CER (%) (LR)	Output (SR)	CER (%) (SR)
GT	majority	0.00	majority	0.00
1 layer	majority	0.00	majority	0.00
2 layers	–	100.00	mnjority	12.50
3 layer	–	100.00	marity	25.00

Table 2: OCR outputs and scores for low-resolution (LR) and super-resolved (SR) versions of text image "majority" across multiple blur layers.

As shown in Figure 4 and Figure 5, the application of the super-resolution model noticeably improves the readability of degraded text images. This is evident from both the qualitative visual results and the quantitative OCR scores presented in the Table 2. Standard OCR systems fail to recognize text in images covered by two or three layers of material, rendering them unreadable in an automated setting. Although the SR model does not always reconstruct every character perfectly, it still yields a substantial improvement in OCR performance. This demonstrates the potential

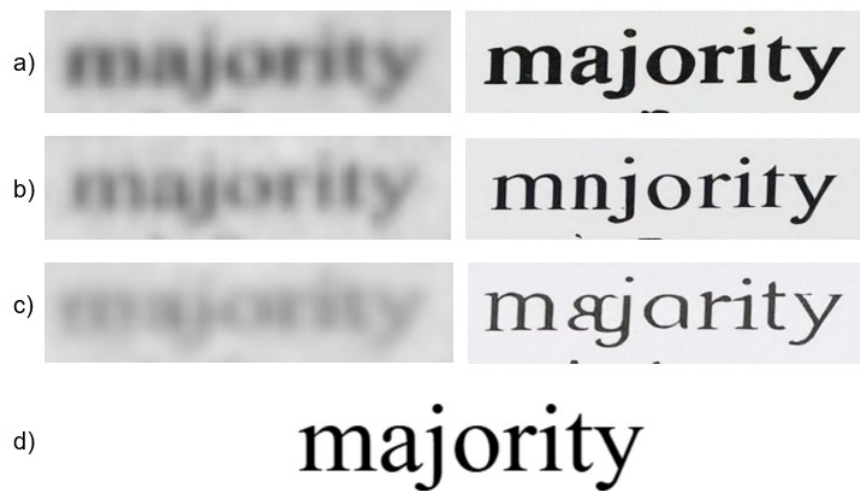


Figure 4: Fabricated samples and the resulting super-resolved results for text covered by a) 1 layer of paper, b) 2 layers of paper, and c) 3 layers of paper. The ground truth is shown in (d).

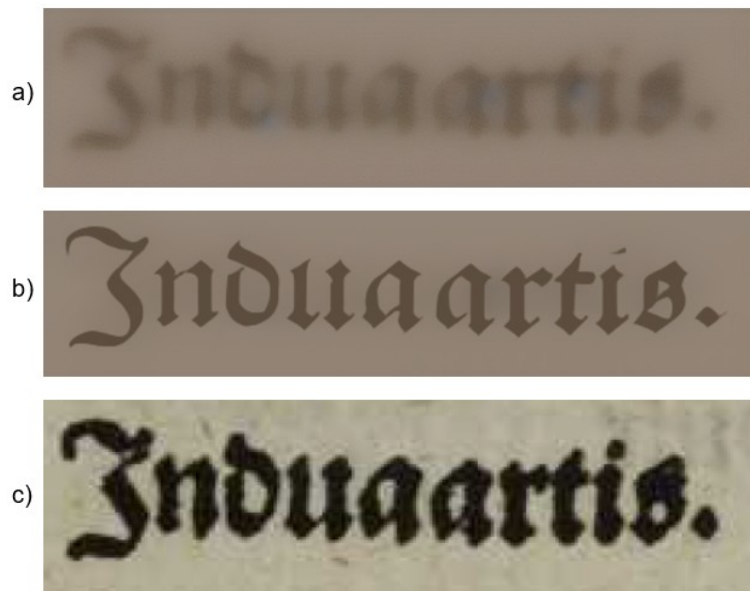


Figure 5: Qualitative performance of model on artificially degraded text "Induaartis". a) The artificially degraded text image, b) the super-resolved result, and c) the ground truth.

of such models to support text recovery in challenging historical materials where conventional methods fall short.

5 Restoring real fragments

5.1 Herlufsholm 24.1

The HSI system used to recover the hidden fragment in Herlufsholm 24.1 includes a line-scan hyperspectral camera [18] with a built-in spectrograph, featuring an IMX990 (1.3 MP) sensor and a Kowa lens with a 35 mm focal length. This system captures 900 spectral channels across the visible to short-wave infrared range (Vis-SWIR, 430 nm to 1700 nm), with a spatial resolution of 1296 pixels along the Y-axis. The camera is mounted 48 cm above a 34 cm-wide conveyor belt

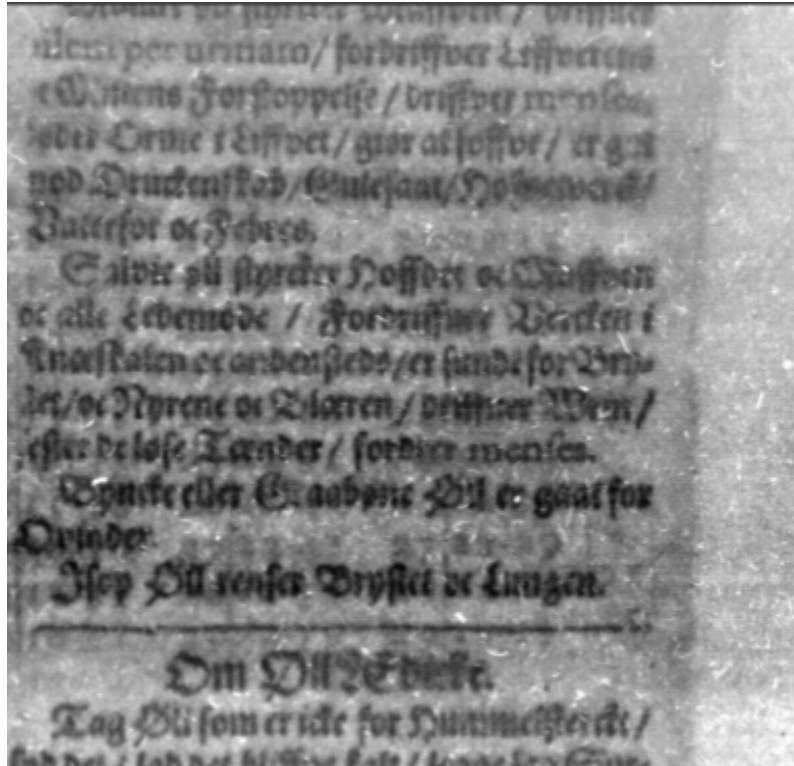


Figure 6: The recovered text image of the fragment embedded in the inner binding of Herlufsholm 24.1.

and operates at a scanning speed of 20 mm/s.

To extract the most readable text image from the hyperspectral data cube, the 900 spectral bands were reduced to four using Principal Component Analysis (PCA). The first four principal components retained over 95% of the total variance. These components were then combined in a way that maximized Shannon entropy [21], enhancing the visibility and contrast of the obscured text. Figure 6 shows a snippet of the recovered fragment on Herlufsholm 24.1 [20].

The recovered text of Herlufsholm 24.1 remains significantly degraded in certain areas. OCR systems developed specifically for historical documents, such as Transkribus [11] and Kraken [14], fail to recognize portions of the text. In some severe cases, the degradation is such that even expert human readers are unable to interpret the content. To assess the effectiveness of the super-resolution model in such scenarios, a number of degraded text images were extracted from the recovered fragment of Herlufsholm 24.1 and used for testing. Figure 7 and Figure 8 illustrate the model’s qualitative performance.

For quantitative evaluation, the OCR score was measured on LR inputs as well as the SR outputs. Similar to Section 4, CER was used as the evaluation metric. Table 3 summarizes the results.

GT	Output (LR)	Output (SR)	CER (%) (LR)	CER (%) (SR)
renser	–	renscr	100.00	16.66
Lungen	–	Lungcn	100.00	16.66
Om	–	Om	100.00	0.00

Table 3: Character Error Rate (CER) for low-resolution (LR) and super-resolved (SR) text predictions from Herlufsholm 24.1.

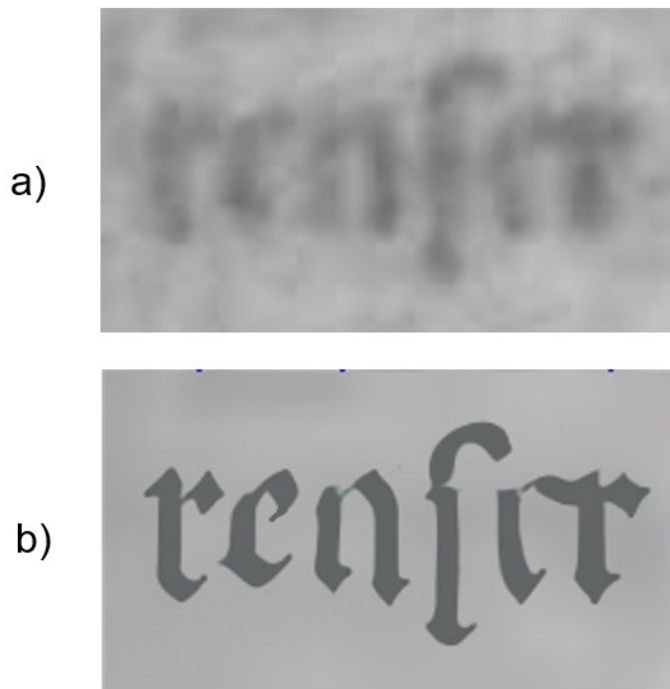


Figure 7: Qualitative performance of the model on degraded text image "renser" extracted from the fragment on Herlufsholm 24.1. The figure shows a) the LR input and b) the SR output.



Figure 8: Qualitative performance of the model on degraded text images "Om" and "Lungen" extracted from the fragment on Herlufsholm 24.1. The figure shows a) the LR inputs and b) the SR outputs.

5.2 Herlufsholm 250.8

Hyperspectral imaging (HSI) was not successful in recovering the hidden text on Herlufsholm 250.8. Preliminary experiments, however, revealed that some text becomes visible under ultraviolet (UV) illumination.

For this purpose, a Canon EOS M50 camera with a 24.1 MP APS-C sensor was used to capture the images. The manuscript cover was illuminated using ultraviolet (UV) light bulbs during imaging. Figure 9 shows the recovered text on Herlufsholm 250.8 as it appears under UV illumination.

The illuminated text on Herlufsholm 250.8 is largely readable to expert human readers. However, due to interference from the complex background, standard OCR models (designed for his-

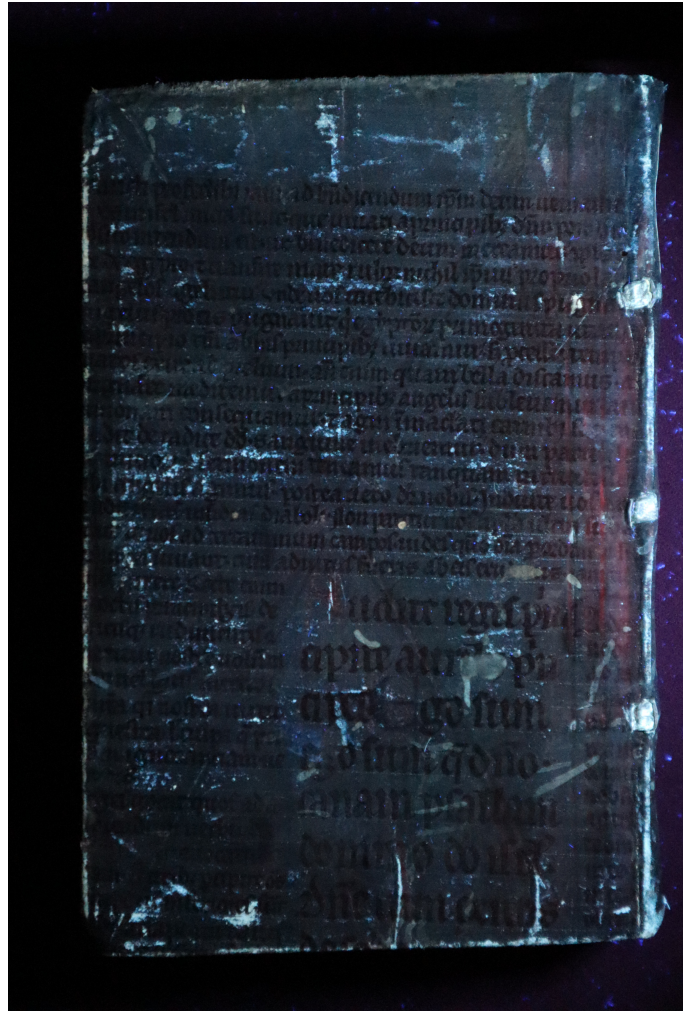


Figure 9: Back cover of Herlufsholm 250.8 under Ultra-violet (UV) light.

torical texts) struggle to accurately recognize the characters. Super-resolution techniques can assist in improving the readability of such degraded images. A qualitative evaluation of extracted text samples is shown in Figure 10 and Figure 11.

As with Herlufsholm 24.1, a quantitative evaluation was also conducted using OCR performance. The CER is used as the primary metric to assess the model’s ability to enhance machine-readability of the recovered text (Table 4).

GT	Output (LR)	Output (SR)	CER (%) (LR)	CER (%) (SR)
nam	an	nam	66.66	0.00
pu	p	pu	50.00	0.00

Table 4: Character Error Rate (CER) for low-resolution (LR) and super-resolved (SR) text predictions on Herlufsholm 250.8.

It can be concluded that, in both the Herlufsholm 250.8 and 24.1 cases, the readability of the extracted text images was improved, based on both qualitative and quantitative OCR metrics. In addition to enhancing legibility, the SR model was able to preserve the stylistic characteristics of the original writings, which is particularly important for historical analysis. These findings demonstrate the potential of SR methods to support the recovery and interpretation of degraded or

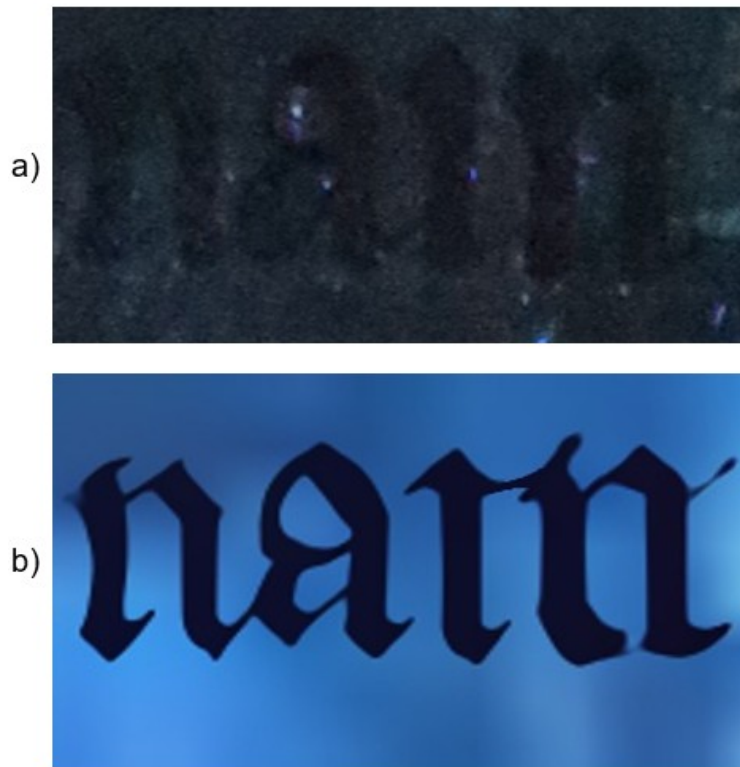


Figure 10: Qualitative performance of the model on degraded text image "nam" extracted from the UV illuminated back-cover of Herlufsholm 250.8. The figure shows a) the LR input and b) the SR output.

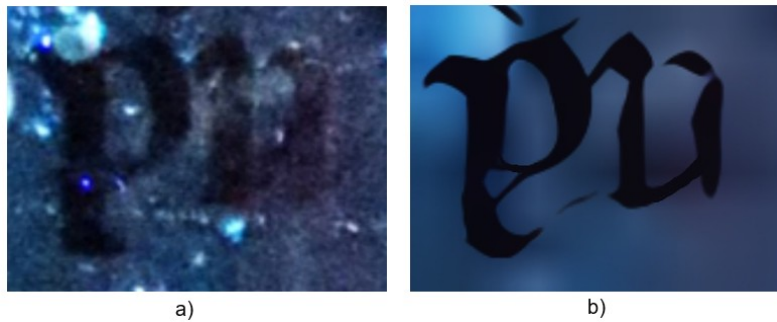


Figure 11: Qualitative performance of the model on degraded text image "pu" extracted from the UV illuminated back-cover of Herlufsholm 250.8. The figure shows a) the LR input and b) the SR output.

obscured historical fragments. However, further evaluation is needed to fully assess the reliability and limitations of such models across different types of degradation and manuscript materials.

6 Discussion and future work

This study demonstrates the potential of blind super-resolution models to enhance the readability of degraded Latin characters in historical manuscript fragments. The application of MARCONet model trained on Latin text styles shows improved readability on both synthetic degradations and

real historical samples. In particular, our qualitative and quantitative analyses suggest that super-resolution can support the deciphering of text in cases where OCR models encounter limitations. The effectiveness of this method on real fragments was tested on only a few samples due to the limited amount of available data. However, further evaluation is required. A more comprehensive assessment of the model’s performance on real fragments will require larger datasets of text images extracted from actual manuscript fragments, as well as expanded collections of controlled, fabricated samples for evaluation and benchmarking purposes.

The successful use of HSI and UV illumination to recover obscured text in Herlufsholm 24.1 and 250.8 also underlines the importance of integrating imaging and computational enhancement techniques. While HSI provided limited benefit in the case of 250.8, UV illumination was effective in revealing text not visible under standard lighting conditions. This reinforces the value of tailored imaging strategies based on material and degradation type.

While some portions of the recovered text image may still be readable by human experts, OCR models typically struggle with severe degradation, which limits the possibility of fully automated transcription. In this context, super-resolution methods can act as a bridge, improving OCR accuracy and enabling more efficient processing. Furthermore, by enhancing visual clarity, such models can also facilitate engagement with historical texts among non-experts, including the general public, who may find degraded fragments inaccessible without computational assistance.

The restoration techniques presented here have the potential to be applied to other items featuring similarly obscured print and manuscript fragments, for example, a group of black fragment bindings similar to Herlufsholm 250.8 held at the University Library of Graz [4]. This opens up new directions for comparative studies and wider deployment of computational restoration tools in manuscript collections beyond the current case study.

A critical consideration for future work is the dependence of blind super-resolution performance on the quality and representativeness of the training data. Since ground truth images of historical fragments are rarely available, artificial degradation pipelines are often used to generate low-resolution samples from high-resolution text images. However, the effectiveness of the blind super-resolution model is closely tied to how well these synthetic degradations approximate real-world degradation. Prior works [24; 25; 26] have proposed diverse degradation models that simulate blur, noise, and compression artifacts, but further investigation is needed to determine which of these models best capture the complexity of historical fragment degradation. Future directions include systematically evaluating different degradation pipelines, fine-tuning the model on fragment-specific styles, and integrating feedback from historical experts into the training loop. These steps will help tailor the model more effectively to the needs of manuscript restoration and cultural heritage preservation.

Acknowledgements

This work is part of the project “*AntCom: From Antiquity to Community*.” The AntCom project has received funding from the European Union’s Horizon 2021 research and innovation programme under the Marie Skłodowska-Curie grant agreement No 101073543.

References

- [1] Camba, A., Gau, M., Hollaus, F., Fiel, S., and Sablatnig, R. “Multispectral Imaging, Image Enhancement, and Automated Writer Identification in Historical Manuscripts”. In: *Manuscript Cultures 7* (2015), pp. 83–91.
- [2] Chen, Jingye, Li, Bin, and Xue, Xiangyang. “Scene Text Telescope: Text-Focused Scene Image Super-Resolution”. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. 2021, pp. 12013–12022.

- [3] Easton, R.L., Knox, K.T., and Christens-Barry, W.A. “Multispectral imaging of the Archimedes palimpsest”. In: *32nd Applied Imagery Pattern Recognition Workshop, 2003. Proceedings.* IEEE. 2003, pp. 111–116. DOI: 10.1109/AIPR.2003.1284258.
- [4] Fehringer, Michaela and Hohl, Werner. “Die Buchbinderwerkstätten der steirischen Klöster. Übersicht der Werkstätten, der verwendeten Prägestempel und der erhaltenen Einbände in der UB Graz”. Tech. rep. Universitätsbibliothek Graz, 2022.
- [5] Fischer, C. and Kakoulli, I. “Multispectral and hyperspectral imaging technologies in conservation: current research and potential applications”. In: *Studies in Conservation* 51, no. sup1 (2006), pp. 3–16.
- [6] Goodfellow, Ian, Pouget-Abadie, Jean, Mirza, Mehdi, Xu, Bing, Warde-Farley, David, Ozair, Sherjil, Courville, Aaron, and Bengio, Yoshua. “Generative Adversarial Nets”. In: *Advances in Neural Information Processing Systems (NeurIPS)*. Vol. 27. 2014, pp. 2672–2680.
- [7] Guan, Shuhao, Lin, Moule, Xu, Cheng, Liu, Xinyi, Zhao, Jinman, Fan, Jiexin, Xu, Qi, and Greene, Derek. “PreP-OCR: A Complete Pipeline for Document Image Restoration and Enhanced OCR Accuracy”. In: *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*. Vienna, Austria: Association for Computational Linguistics, 2025, pp. 15413–15425. DOI: 10.18653/v1/2025.acl-long.749.
- [8] Holck, Jakob Povl. *Den gamle verdens magi: Bogsamlingen fra Herlufsholm på Syddansk Universitetsbibliotek*. Syddansk Universitetsbibliotek, 2015.
- [9] Holck, Jakob Povl. “Herlufsholm-samlingen på Syddansk Universitetsbibliotek og forskningspotentialet - om forskellige fund af fragmenter”. In: *Studier i Nordisk*, no. 2016-2018 (2022), pp. 5–40.
- [10] Holck, Jakob Povl. “Nye teknologier og Herlufsholm-samlingen: Det skjulte bibliotek kommer for en dag”. In: *Bibliotekshistorie* 13 (2018), pp. 38–66.
- [11] Kahle, Peter, Colutto, Simon, Hackl, Günter, and Mühlberger, Günter. “Transkribus - A Service Platform for Transcription, Recognition and Retrieval of Historical Documents”. In: *14th IAPR International Conference on Document Analysis and Recognition (ICDAR)*. Vol. 4. 2017, pp. 19–24. DOI: 10.1109/ICDAR.2017.307.
- [12] Karras, Tero, Laine, Samuli, and Aila, Timo. “A Style-Based Generator Architecture for Generative Adversarial Networks”. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. 2019, pp. 4401–4410.
- [13] Karras, Tero, Laine, Samuli, Aittala, Miika, Hellsten, Janne, Lehtinen, Jaakko, and Aila, Timo. “Analyzing and Improving the Image Quality of StyleGAN”. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. 2020, pp. 8110–8119.
- [14] Kiessling, Benjamin. “Kraken: A Turn-Key OCR System Optimized for Historical and Non-Latin Script Material”. <https://kraken.re/>. Version 6.0.0 (Sept 3 2025). Apache-2.0 license. 2019.
- [15] Kingma, Diederik P and Ba, Jimmy. “Adam: A Method for Stochastic Optimization”. In: *arXiv preprint arXiv:1412.6980* (2014).
- [16] Li, Xiaoming, Zuo, Wangmeng, and Loy, Chen Change. “Learning Generative Structure Prior for Blind Text Image Super-resolution”. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. 2023.

- [17] Ma, Jianqi, Liang, Zhetong, and Zhang, Lei. “A Text Attention Network for Spatial Deformation Robust Scene Text Image Super-Resolution”. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. 2022, pp. 19009–19018.
- [18] Newtec M.G. “Pushbroom Hyperspectral Imaging System Buteo”. <https://www.newtec.com/machines/quality-control/buteo>. Last accessed 2025-03-06. 2025.
- [19] Outzen, S., Christensen, K., and Sydøstdanmark, Museum. *Birgitte Gøye & Herluf Trolles bøger*. Museum Sydøstdanmark, 2015.
- [20] Pourahmadi, Baharan, Andersen, Morten Sielnik, Holck, Jakob Povl, Jensen, Mogens Kragssig, Frandsen, Mads Toudal, and Eriksen, René Lynge. “Challenges in Revealing Readable Text from Fragments Hidden in Book Bindings: A Case Study from the Herlufsholm Collection”. In: *Document Analysis and Recognition – ICDAR 2025*. 2026, pp. 493–508.
- [21] Shannon, Claude E. “A Mathematical Theory of Communication”. In: *Bell System Technical Journal* 27, no. 3 (1948), pp. 379–423. DOI: 10.1002/j.1538-7305.1948.tb00917.x.
- [22] Tahkokallio, Jaakko Kalervo, Niskanen, Samu, and Willoughby, James. “Manuscript Fragments”. In: *Routledge Resources Online–Medieval Studies*. Routledge, 2023.
- [23] Tonazzini, A, Salerno, E, Abdel-Salam, ZA, Harith, MA, Marras, L, Botto, A, and Palleschi, V. “Analytical and mathematical methods for revealing hidden details in ancient manuscripts and paintings: A review”. In: *Journal of Advanced Research* 17 (2019), pp. 31–42. DOI: <https://doi.org/10.1016/j.jare.2019.01.003>.
- [24] Wang, Xintao, Xie, Liangbin, Dong, Chao, and Shan, Ying. “Real-ESRGAN: Training Real-World Blind Super-Resolution with Pure Synthetic Data”. In: *Proceedings of the IEEE/CVF International Conference on Computer Vision Workshops (ICCVW)*. 2021, pp. 1905–1914.
- [25] Zhang, Kai, Liang, Jingyun, Van Gool, Luc, and Timofte, Radu. “Designing a Practical Degradation Model for Deep Blind Image Super-Resolution”. In: *IEEE International Conference on Computer Vision*. 2021, pp. 4791–4800.
- [26] Zhang, Kaihao, Luo, Wenhan, Zhong, Yiran, Ma, Lin, Stenger, Bjorn, Liu, Wei, and Li, Hongdong. “Deblurring by realistic blurring”. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. 2020, pp. 2737–2746.