



ΠΑΝΕΠΙΣΤΗΜΙΟ
ΘΕΣΣΑΛΙΑΣ

ΣΧΟΛΗ ΘΕΤΙΚΩΝ ΕΠΙΣΤΗΜΩΝ

ΤΜΗΜΑ ΠΛΗΡΟΦΟΡΙΚΗΣ ΚΑΙ ΤΗΛΕΠΙΚΟΙΝΩΝΙΩΝ

Ευφυής Αναγνώριση και Αξιολόγηση Φθορών σε Έργα Τέχνης

ΧΡΥΣΑΝΘΗ ΠΕΓΙΟΓΛΟΥ

ΠΤΥΧΙΑΚΗ ΕΡΓΑΣΙΑ

ΥΠΕΥΘΥΝΟΣ

ΚΟΛΟΜΒΑΤΣΟΣ ΚΩΝΣΤΑΝΤΙΝΟΣ
ΑΝΑΠΛΗΡΩΤΗΣ ΚΑΘΗΓΗΤΗΣ

Λαμία έτος 2025



ΠΑΝΕΠΙΣΤΗΜΙΟ
ΘΕΣΣΑΛΙΑΣ

ΣΧΟΛΗ ΘΕΤΙΚΩΝ ΕΠΙΣΤΗΜΩΝ

ΤΜΗΜΑ ΠΛΗΡΟΦΟΡΙΚΗΣ ΚΑΙ ΤΗΛΕΠΙΚΟΙΝΩΝΙΩΝ

Ευφυής Αναγνώριση και Αξιολόγηση Φθορών σε Έργα Τέχνης

ΧΡΥΣΑΝΘΗ ΠΕΓΓΙΟΓΛΟΥ

ΠΤΥΧΙΑΚΗ ΕΡΓΑΣΙΑ

ΥΠΕΥΘΥΝΟΣ

ΚΟΛΟΜΒΑΤΣΟΣ ΚΩΝΣΤΑΝΤΙΝΟΣ
ΑΝΑΠΛΗΡΩΤΗΣ ΚΑΘΗΓΗΤΗΣ

Λαμία έτος 2025



UNIVERSITY OF
THESSALY

SCHOOL OF SCIENCE

DEPARTMENT OF COMPUTER SCIENCE & TELECOMMUNICATIONS

Intelligent Damage Detection and Assessment in Artworks

CHRYSANTHI PEGIOGLOU

FINAL THESIS

ADVISOR

KOLOMVASOS KONSTANTINOS
ASSOCIATE PROFESSOR

Lamia year 2025

«Με ατομική μου ευθύνη και γνωρίζοντας τις κυρώσεις ⁽¹⁾, που προβλέπονται από της διατάξεις της παρ. 6 του άρθρου 22 του Ν. 1599/1986, δηλώνω ότι:

1. Δεν παραθέτω κομμάτια βιβλίων ή άρθρων ή εργασιών άλλων αυτολεξεί **χωρίς να τα περικλείω σε εισαγωγικά** και χωρίς να αναφέρω το συγγραφέα, τη χρονολογία, τη σελίδα. Η αυτολεξεί παράθεση χωρίς εισαγωγικά χωρίς αναφορά στην πηγή, είναι λογοκλοπή. Πέραν της αυτολεξεί παράθεσης, λογοκλοπή θεωρείται και η παράφραση εδαφίων από έργα άλλων, συμπεριλαμβανομένων και έργων συμφοιτητών μου, καθώς και η παράθεση στοιχείων που άλλοι συνέλεξαν ή επεξεργάστηκαν, χωρίς αναφορά στην πηγή. Αναφέρω πάντοτε με πληρότητα την πηγή κάτω από τον πίνακα ή σχέδιο, όπως στα παραθέματα.
2. Δέχομαι ότι η αυτολεξεί **παράθεση χωρίς εισαγωγικά**, ακόμα κι αν συνοδεύεται από αναφορά στην πηγή σε κάποιο άλλο σημείο του κειμένου ή στο τέλος του, είναι αντιγραφή. Η αναφορά στην πηγή στο τέλος π.χ. μιας παραγράφου ή μιας σελίδας, δεν δικαιολογεί συρραφή εδαφίων έργου άλλου συγγραφέα, έστω και παραφρασμένων, και παρουσίασή τους ως δική μου εργασία.
3. Δέχομαι ότι υπάρχει επίσης περιορισμός στο μέγεθος και στη συχνότητα των παραθεμάτων που μπορώ να εντάξω στην εργασία μου εντός εισαγωγικών. Κάθε μεγάλο παράθεμα (π.χ. σε πίνακα ή πλαίσιο, κλπ), προϋποθέτει ειδικές ρυθμίσεις, και όταν δημοσιεύεται προϋποθέτει την άδεια του συγγραφέα ή του εκδότη. Το ίδιο και οι πίνακες και τα σχέδια
4. Δέχομαι όλες τις συνέπειες σε περίπτωση λογοκλοπής ή αντιγραφής.

Ημερομηνία: 27/6/2025

Ο – Η Δηλ.

ΧΡΥΣΑΝΘΗ ΠΕΓΙΟΓΛΟΥ

(1) «Όποιος εν γνώσει του δηλώνει ψευδή γεγονότα ή αρνείται ή αποκρύπτει τα αληθινά με έγγραφη υπεύθυνη δήλωση του άρθρου 8 παρ. 4 Ν. 1599/1986 τιμωρείται με φυλάκιση τουλάχιστον τριών μηνών. Εάν ο υπαίτιος αυτών των πράξεων σκόπευε να προσπορίσει στον εαυτόν του ή σε άλλον περιουσιακό όφελος βλάπτοντας τρίτον ή σκόπευε να βλάψει άλλον, τιμωρείται με κάθειρξη μέχρι 10 ετών.»

ΠΕΡΙΛΗΨΗ

Η μελέτη αυτή παρουσιάζει ένα ολοκληρωμένο πλαίσιο για τον αυτοματοποιημένο εντοπισμό και τη σημασιολογική τμηματοποίηση φυσικών φθορών σε ψηφιοποιημένα πολιτιστικά τεκμήρια, αξιοποιώντας τεχνικές βαθιάς μάθησης και επεξηγήσιμης τεχνητής νοημοσύνης. Χρησιμοποιώντας το σύνολο δεδομένων ARTeFACT, το σύστημα εντοπίζει περιοχές φθοράς — όπως ρωγμές, λεκέδες και απώλεια υλικού — σε ιστορικά έγγραφα και έργα τέχνης.

Το μοντέλο βασίζεται σε αρχιτεκτονική U-Net με βάση ResNet-50 και εκπαιδεύεται σε σύνολο δεδομένων με επισημάνσεις σε επίπεδο εικονοστοιχείου, επιτρέποντας ακριβή δυαδική τμηματοποίηση. Η διαδικασία εκπαίδευσης περιλαμβάνει προηγμένες τεχνικές όπως ενίσχυση δεδομένων, εκπαίδευση με μικτή ακρίβεια και βελτιστοποίηση με Εκθετικό Κινούμενο Μέσο (EMA) για βελτίωση της απόδοσης και της γενίκευσης.

Για τη διασφάλιση διαφάνειας και ερμηνευσιμότητας, το πλαίσιο ενσωματώνει χάρτες σαφήνειας και μετρικές φθοράς (π.χ. επιφάνεια, εκκεντρότητα, συμπαγής μορφή) για οπτική και ποσοτική αξιολόγηση. Ένα διαδραστικό περιβάλλον με χρήση `ipywidgets` επιτρέπει στους χρήστες να εξερευνούν προβλέψεις, πραγματικά δεδομένα και εξηγήσεις του μοντέλου, ενισχύοντας τη συνεργασία ανθρώπου-μηχανής.

ABSTRACT

This study presents a comprehensive framework for the automated detection and semantic segmentation of physical damage in digitized cultural heritage artifacts, leveraging deep learning and explainable AI methodologies. Using the ARTeFACT dataset, the system identifies damaged areas—such as cracks, stains, and material loss—on historical documents and artworks.

The model is based on a U-Net architecture with a ResNet-50 backbone and is trained on a pixel-level annotated dataset, enabling precise binary segmentation. The training pipeline incorporates advanced techniques such as data augmentation, mixed precision training, and Exponential Moving Average (EMA) optimization to improve performance and generalization.

To ensure transparency and interpretability, the framework integrates saliency maps and damage metrics (e.g., area, eccentricity, solidity) for both visual and quantitative assessment. An interactive interface built with ipywidgets supports human-in-the-loop collaboration by allowing users to explore predictions, ground truth, and model explanations.

Πίνακας Περιεχομένων

ΠΕΡΙΛΗΨΗ	I
ABSTRACT	III
ΚΕΦΑΛΑΙΟ 1 ΕΙΣΑΓΩΓΗ.....	2
1.1. Η ΣΗΜΑΣΙΑ ΤΗΣ ΑΝΙΧΝΕΥΣΗΣ ΚΑΙ ΑΞΙΟΛΟΓΗΣΗΣ ΦΘΟΡΩΝ ΣΤΗΝ ΤΕΧΝΗ	2
1.2. ΕΞΕΛΙΞΗ ΤΗΣ ΤΕΧΝΗΤΗΣ ΝΟΗΜΟΣΥΝΗΣ ΣΤΗ ΣΥΝΤΗΡΗΣΗ ΈΡΓΩΝ ΤΕΧΝΗΣ	2
1.3. ΕΡΕΥΝΗΤΙΚΑ ΟΡΟΣΗΜΑ ΣΤΗΝ ΑΥΤΟΜΑΤΗ ΑΝΙΧΝΕΥΣΗ ΦΘΟΡΩΝ	3
1.4. ΣΤΟΧΟΙ, ΕΡΕΥΝΗΤΙΚΑ ΕΡΩΤΗΜΑΤΑ ΚΑΙ ΣΥΜΒΟΛΗ ΤΗΣ ΠΑΡΟΥΣΑΣ ΕΡΓΑΣΙΑΣ	3
ΚΕΦΑΛΑΙΟ 2 ΒΙΒΛΙΟΓΡΑΦΙΚΗ ΕΠΙΣΚΟΠΗΣΗ	6
2.1. ΕΙΣΑΓΩΓΗ	6
2.1.A. ΕΠΕΞΗΓΗΣΗ ΤΕΧΝΙΚΩΝ ΚΑΙ ΑΛΓΟΡΙΘΜΩΝ.....	6
2.2. ΚΡΙΤΙΚΗ ΑΠΟΤΙΜΗΣΗ ΤΕΧΝΙΚΩΝ ΚΑΙ ΑΠΟΤΕΛΕΣΜΑΤΩΝ	7
2.2.A. ΑΝΙΧΝΕΥΣΗ ΦΘΟΡΩΝ ΜΕ ΠΟΛΥΤΡΟΠΙΚΑ ΔΕΔΟΜΕΝΑ ΚΑΙ CNNs	7
2.2.B. ΑΠΟΚΑΤΑΣΤΑΣΗ ΕΙΚΟΝΑΣ ΚΑΙ ΑΝΟΙΧΤΑ ΕΡΓΑΛΕΙΑ	8
2.2.G. ΕΠΙΣΗΜΕΙΩΣΗ, ΕΞΕΙΔΙΚΕΥΣΗ ΚΑΙ ΕΠΕΚΤΑΣΗ ΔΕΔΟΜΕΝΩΝ	9
2.2.D. ΕΝΣΩΜΑΤΩΜΕΝΑ ΕΡΓΑΛΕΙΑ ΚΑΙ ΠΡΟΛΗΠΤΙΚΗ ΣΥΝΤΗΡΗΣΗ	10
2.2.E. ΗΘΙΚΕΣ ΚΑΙ ΔΙΕΠΙΣΤΗΜΟΝΙΚΕΣ ΠΡΟΕΚΤΑΣΕΙΣ	11
2.3. ΕΡΜΗΝΕΙΑ ΚΑΙ ΕΦΑΡΜΟΣΙΜΟΤΗΤΑ ΤΩΝ ΕΥΡΗΜΑΤΩΝ	12
2.3.A. ΠΡΑΚΤΙΚΗ ΑΞΙΑ ΚΑΙ ΥΠΟΣΤΗΡΙΞΗ ΑΠΟΦΑΣΗΣ	12
2.3.B. ΕΡΜΗΝΕΥΣΙΜΟΤΗΤΑ ΚΑΙ ΕΜΠΙΣΤΟΣΥΝΗ	12
2.3.G. ΓΕΝΙΚΕΥΣΗ ΚΑΙ ΑΝΘΕΚΤΙΚΟΤΗΤΑ ΜΟΝΤΕΛΩΝ	12
2.3.D. ΣΥΓΚΡΙΤΙΚΗ ΑΝΑΛΥΣΗ ΕΠΙΛΕΓΜΕΝΩΝ ΜΕΛΕΤΩΝ	12
2.3.E. ΕΝΣΩΜΑΤΩΣΗ ΣΕ ΡΟΕΣ ΕΡΓΑΣΙΑΣ ΚΑΙ ΑΝΘΡΩΠΙΝΗ ΣΥΝΕΡΓΑΣΙΑ	15
2.4. ΘΕΩΡΗΤΙΚΕΣ, ΜΕΘΟΔΟΛΟΓΙΚΕΣ ΚΑΙ ΤΕΧΝΙΚΕΣ ΠΡΟΚΛΗΣΕΙΣ	15
2.4.A. ΠΕΡΙΟΡΙΣΜΟΙ ΤΩΝ ΥΦΙΣΤΑΜΕΝΩΝ ΜΕΘΟΔΩΝ	15
2.4.B. ΜΕΘΟΔΟΛΟΓΙΚΕΣ ΠΡΟΤΑΣΕΙΣ ΓΙΑ ΜΕΛΛΟΝΤΙΚΗ ΈΡΕΥΝΑ	16
2.5. ΚΟΙΝΩΝΙΚΕΣ, ΗΘΙΚΕΣ ΚΑΙ ΘΕΣΜΙΚΕΣ ΠΤΥΧΕΣ.....	16
2.6. ΠΡΟΤΑΣΕΙΣ ΓΙΑ ΜΕΛΛΟΝΤΙΚΗ ΈΡΕΥΝΑ ΚΑΙ ΕΦΑΡΜΟΓΗ	17
2.7. ΣΥΝΟΛΙΚΗ ΑΠΟΤΙΜΗΣΗ ΚΑΙ ΣΥΝΔΕΣΗ ΜΕ ΤΟΥΣ ΣΤΟΧΟΥΣ ΤΗΣ ΕΡΓΑΣΙΑΣ	17
ΚΕΦΑΛΑΙΟ 3 ΜΕΘΟΔΟΛΟΓΙΑ	19
3.1. ΣΧΕΔΙΑΣΜΟΣ ΤΗΣ ΈΡΕΥΝΑΣ ΚΑΙ ΜΕΘΟΔΟΛΟΓΙΚΗ ΤΕΚΜΗΡΙΩΣΗ.....	19
3.2. ΠΕΡΙΓΡΑΦΗ ΤΟΥ ΣΥΝΟΛΟΥ ΔΕΔΟΜΕΝΩΝ: ARTEFACT	19
3.3. ΑΠΟ ΤΗΝ ΠΟΛΥΚΑΤΗΓΟΡΙΚΗ ΣΤΗΝ ΔΥΑΔΙΚΗ ΣΗΜΑΣΙΟΛΟΓΙΚΗ ΤΜΗΜΑΤΟΠΟΙΗΣΗ (BINARY SEMANTIC SEGMENTATION): ΜΙΑ ΕΞΕΛΙΚΤΙΚΗ ΠΟΡΕΙΑ	20
3.3.A. ΕΠΙΒΕΒΑΙΩΣΗ ΑΠΟ ΤΗ ΒΙΒΛΙΟΓΡΑΦΙΑ: ΠΕΡΙΟΡΙΣΜΟΙ ΤΗΣ ΤΜΗΜΑΤΟΠΟΙΗΣΗΣ ΣΤΟ ARTEFACT	22
3.4. ΤΕΛΙΚΗ ΠΡΟΣΕΓΓΙΣΗ: ΔΥΑΔΙΚΗ ΣΗΜΑΣΙΟΛΟΓΙΚΗ ΤΜΗΜΑΤΟΠΟΙΗΣΗ (BINARY SEMANTIC SEGMENTATION) ΜΕ U-NET	22
3.4.A. ΔΗΜΙΟΥΡΓΙΑ ΕΤΙΚΕΤΩΝ	23
3.4.B. ΠΡΟΕΠΕΞΕΡΓΑΣΙΑ ΚΑΙ ΕΜΠΛΟΥΤΙΣΜΟΣ ΔΕΔΟΜΕΝΩΝ (DATA PREPROCESSING & AUGMENTATION)	23
3.4.G. ΕΚΤΙΜΗΣΗ ΣΟΒΑΡΟΤΗΤΑΣ ΦΘΟΡΑΣ (DAMAGE SEVERITY ESTIMATION)	24

3.5. ΕΡΜΗΝΕΥΣΙΜΟΤΗΤΑ ΚΑΙ ΣΥΝΕΡΓΑΣΙΑ ΑΝΘΡΩΠΟΥ–ΜΗΧΑΝΗΣ (EXPLAINABILITY AND HUMAN-IN-THE-LOOP INTERACTION)	24
3.5.A. ΟΠΤΙΚΟΠΟΙΗΣΗ ΠΡΟΣΟΧΗΣ ΜΕΣΩ SALIENCY MAPS	25
3.5.B. ΔΙΑΔΡΑΣΤΙΚΗ ΔΙΕΠΑΦΗ ΜΕ ΕΡΜΗΝΕΥΣΙΜΟΤΗΤΑ (INTERACTIVE EXPLAINABLE INTERFACE)	25
3.5.Γ. ΕΝΑΛΛΑΚΤΙΚΕΣ ΤΕΧΝΙΚΕΣ EXPLAINABLE AI	26
3.5.Δ. ΠΡΟΣ'ΕΝΑ ΕΠΕΞΗΓΗΣΙΜΟ ΚΑΙ ΠΡΟΣΑΡΜΟΣΙΜΟ ΣΥΣΤΗΜΑ	26
3.6. ΑΞΙΟΛΟΓΗΣΗ ΓΕΝΙΚΕΥΣΗΣ ΚΑΙ ΠΕΙΡΑΜΑΤΙΚΗ ΕΠΕΚΤΑΣΗ (GENERALIZATION ASSESSMENT AND EXPERIMENTAL EXTENSIONS)	26
3.6.A. ΤΡΕΧΟΥΣΑ ΑΞΙΟΛΟΓΗΣΗ ΚΑΙ ΠΡΟΤΕΙΝΟΜΕΝΕΣ ΕΠΕΚΤΑΣΕΙΣ ΜΕ ΜΕΤΑΔΕΔΟΜΕΝΑ	27
3.7. ΤΕΧΝΙΚΗ ΠΕΡΙΓΡΑΦΗ ΤΗΣ ΥΛΟΠΟΙΗΣΗΣ	27
3.7.A. ΔΟΜΗ ΣΥΣΤΗΜΑΤΟΣ ΚΑΙ ΡΟΗ ΕΡΓΑΣΙΑΣ	28
3.7.B. ΕΡΓΑΛΕΙΑ ΚΑΙ ΒΙΒΛΙΟΘΗΚΕΣ	30
3.8. ΑΝΑΜΕΝΟΜΕΝΕΣ ΠΡΟΚΛΗΣΕΙΣ ΚΑΙ ΤΕΧΝΙΚΕΣ ΠΑΡΑΔΟΧΕΣ	30
3.8.A. ΑΝΙΣΟΡΡΟΠΙΑ ΚΑΙ ΕΤΕΡΟΓΕΝΕΙΑ ΔΕΔΟΜΕΝΩΝ	30
3.8.B. ΠΕΡΙΟΡΙΣΜΟΙ ΥΠΟΛΟΓΙΣΤΙΚΩΝ ΠΟΡΩΝ	30
3.8.Γ. ΕΡΜΗΝΕΥΣΙΜΟΤΗΤΑ ΚΑΙ ΑΝΘΡΩΠΙΝΗ ΕΜΠΙΣΤΟΣΥΝΗ	31
3.8.Δ. ΓΕΝΙΚΕΥΣΗ ΚΑΙ ΕΦΑΡΜΟΓΗ ΣΕ ΝΕΑ ΠΕΡΙΒΑΛΛΟΝΤΑ	31
3.8.E. ΠΑΡΑΔΟΧΕΣ ΚΑΙ ΠΕΡΙΟΡΙΣΜΟΙ ΤΗΣ ΥΛΟΠΟΙΗΣΗΣ	31
 ΚΕΦΑΛΑΙΟ 4 ΑΝΑΛΥΣΗ ΚΑΙ ΣΥΖΗΤΗΣΗ ΑΠΟΤΕΛΕΣΜΑΤΩΝ	 32
4.1. ΕΠΙΣΚΟΠΗΣΗ	32
4.2. ΡΟΗ ΠΡΟΓΡΑΜΜΑΤΟΣ ΚΑΙ ΠΑΡΑΜΕΤΡΟΠΟΙΗΣΗ	34
4.2.A. ΡΟΗ ΕΠΕΞΕΡΓΑΣΙΑΣ ΚΑΙ ΕΚΠΑΙΔΕΥΣΗΣ	34
4.2.B. ΠΑΡΑΜΕΤΡΟΠΟΙΗΣΗ ΤΟΥ ΣΥΣΤΗΜΑΤΟΣ	35
4.3. ΑΠΟΔΟΣΗ ΕΚΠΑΙΔΕΥΣΗΣ ΚΑΙ ΕΠΙΚΥΡΩΣΗΣ	35
4.4. ΑΞΙΟΛΟΓΗΣΗ ΣΤΟ TEST SET	36
4.5. ΟΠΤΙΚΟΠΟΙΗΣΗ ΠΡΟΒΛΕΨΕΩΝ	37
4.6. ΑΝΑΛΥΣΗ ΦΘΟΡΑΣ ΚΑΙ ΜΕΤΡΙΚΕΣ	39
4.7 ΕΡΜΗΝΕΥΣΙΜΟΤΗΤΑ ΚΑΙ ΔΙΑΔΡΑΣΤΙΚΟΤΗΤΑ	40
4.8. ΣΥΓΚΡΙΤΙΚΗ ΑΝΑΛΥΣΗ ΚΑΙ ΠΕΡΙΟΡΙΣΜΟΙ	41
4.9. ΣΥΝΟΨΗ	42
 ΚΕΦΑΛΑΙΟ 5 ΣΥΜΠΕΡΑΣΜΑΤΑ ΚΑΙ ΜΕΛΛΟΝΤΙΚΗ ΕΡΓΑΣΙΑ	 44
5.1 ΣΥΜΠΕΡΑΣΜΑΤΑ	44
5.2. ΜΕΛΛΟΝΤΙΚΗ ΕΡΓΑΣΙΑ	44
 ΒΙΒΛΙΟΓΡΑΦΙΑ	 46
 ΠΑΡΑΡΤΗΜΑ Α - ΠΙΝΑΚΑΣ ΌΡΩΝ ΚΑΙ ΣΥΝΤΟΜΟΓΡΑΦΙΩΝ	 49
A. ΤΕΧΝΙΚΟΙ ΌΡΟΙ ΚΑΙ ΜΕΤΡΙΚΕΣ ΑΠΟΔΟΣΗΣ (TECHNICAL TERMS & PERFORMANCE METRICS)	49
B. ΌΡΟΙ ΣΥΝΤΗΡΗΣΗΣ ΚΑΙ ΤΕΧΝΟΛΟΓΙΑΣ ΤΕΧΝΗΣ (ART CONSERVATION TERMS)	50

ΚΕΦΑΛΑΙΟ 1 Εισαγωγή

1.1. Η σημασία της ανίχνευσης και αξιολόγησης φθορών στην τέχνη

Η διατήρηση και αποκατάσταση ζωγραφικών έργων τέχνης αποτελεί διαχρονικό στόχο της επιστημονικής και πολιτιστικής κοινότητας, καθώς τα έργα σε καμβά είναι ευάλωτα σε φθορές που προκαλούνται από τη γήρανση, περιβαλλοντικούς παράγοντες και ανθρώπινες παρεμβάσεις (Sizyakin, et al., 2018; Garcia-Moreno, et al., 2024a; Huang, et al., 2020). Η ακριβής ανίχνευση και αξιολόγηση των ζημιών είναι ζωτικής σημασίας για τη λήψη τεκμηριωμένων αποφάσεων από τους συντηρητές. Ωστόσο, οι παραδοσιακές μέθοδοι βασίζονται σε χρονοβόρες, υποκειμενικές και συχνά μη επαναλήψιμες διαδικασίες, γεγονός που καθιστά την εφαρμογή τους ιδιαίτερα απαιτητική, ιδίως σε μεγάλες συλλογές ή έργα με σύνθετες τεχνικές (Huang, et al., 2020; Garcia-Moreno, et al., 2024a; Dulecha, et al., 2019).

Τα τελευταία χρόνια, η ραγδαία πρόοδος στον τομέα της μηχανικής μάθησης (ML) και της βαθιάς μάθησης (DL) έχει προσφέρει νέα εργαλεία για την αυτοματοποιημένη ανίχνευση και αξιολόγηση φθορών σε ψηφιακές εικόνες έργων τέχνης, ενισχύοντας την αντικειμενικότητα, την ταχύτητα και την ακρίβεια της διαδικασίας (Sizyakin, et al., 2018; Sizyakin, et al., 2022; Garcia-Moreno, et al., 2024a).

1.2. Εξέλιξη της τεχνητής νοημοσύνης στη συντήρηση έργων τέχνης

Οι πρόσφατες εξελίξεις στην τεχνητή νοημοσύνη έχουν οδηγήσει σε σημαντικές καινοτομίες στον τομέα της ανάλυσης και αποκατάστασης έργων τέχνης. Η αξιοποίηση τεχνικών όπως τα συνελκτικά νευρωνικά δίκτυα (CNNs), τα γενετικά ανταγωνιστικά δίκτυα (GANs), τα diffusion models, καθώς και μεθόδων ενεργής μάθησης και πολυτροπικής ανάλυσης δεδομένων, έχει επιτρέψει την αυτόματη ανίχνευση ρωγμών, απώλειας χρώματος, μούχλας και άλλων τύπων φθοράς με ακρίβεια που συχνά προσεγγίζει ή και υπερβαίνει τις παραδοσιακές μεθόδους (Huang, et al., 2020; Sizyakin, et al., 2022; Nordin, et al., 2025; Kumar & Gupta, 2025).

Η ανάπτυξη εξειδικευμένων συνόλων δεδομένων με επισημάνσεις από ειδικούς συντηρητές, σε συνδυασμό με την αξιοποίηση πολυτροπικών δεδομένων, έχει ενισχύσει σημαντικά τη δυνατότητα γενίκευσης και την αξιοπιστία των μοντέλων τεχνητής νοημοσύνης (Huang, et al., 2020; Garcia-Moreno, et al., 2024a; Ivanova, et al., 2024; Ivanova, et al., 2025). Επιπλέον, η ενσωμάτωση εργαλείων επεξηγήσιμης τεχνητής νοημοσύνης (explainable AI), όπως οι saliency maps και το Grad-CAM, καθώς και η προσέγγιση συνεργασίας ανθρώπου-μηχανής (human-in-the-loop), ενισχύουν την αποδοχή και την πρακτική εφαρμογή αυτών των τεχνολογιών από την κοινότητα των συντηρητών (Basu, et al., 2023; Fontaine, 2024).

1.3. Ερευνητικά ορόσημα στην αυτόματη ανίχνευση φθορών

Σημαντικά ερευνητικά ορόσημα περιλαμβάνουν την εφαρμογή τεχνικών βαθιάς μάθησης για την ανίχνευση ρωγμών μέσω αρχιτεκτονικών όπως τα U-Net και Mask R-CNN (Sizyakin, et al., 2018; Sizyakin, et al., 2022) , την εικονική αποκατάσταση (virtual inpainting) με GANs και diffusion models , καθώς και την ανάλυση φθορών όπως η μούχλα και οι lacunae με σύγχρονες τεχνικές εξαγωγής χαρακτηριστικών και ταξινόμησης (Nordin, et al., 2025; Garcia-Moreno, et al., 2024a). Οι μελέτες αυτές καταδεικνύουν τη δυναμική των αλγορίθμων μηχανικής μάθησης στην υποστήριξη της λήψης αποφάσεων και στην αυτοματοποίηση κρίσιμων διαδικασιών συντήρησης.

Πέραν της ανίχνευσης, η αυτόματη αξιολόγηση της σοβαρότητας και του τύπου της φθοράς αποτελεί κρίσιμο βήμα για τον σχεδιασμό στοχευμένων και αποτελεσματικών παρεμβάσεων. Πρόσφατες έρευνες έχουν επικεντρωθεί στην ανάπτυξη μεθόδων που όχι μόνο εντοπίζουν τις φθορές, αλλά τις κατηγοριοποιούν και εκτιμούν το επίπεδο επικινδυνότητάς τους (Califano, et al., 2022; Nordin, et al., 2025; Huang, et al., 2020). Για παράδειγμα, έχουν προταθεί συστήματα που διακρίνουν μεταξύ επιφανειακών και δομικών ρωγμών ή αξιολογούν την έκταση της απώλειας χρώματος και τη δυνητική της επίδραση στη σταθερότητα του έργου (Califano, et al., 2022; Huang, et al., 2020). Επιπλέον, μοντέλα όπως το XGBoost έχουν χρησιμοποιηθεί για την πρόβλεψη της εξέλιξης φθορών και την εκτίμηση του κινδύνου δομικής αστοχίας (Califano, et al., 2022).

Παρά τη σημαντική πρόοδο, η βιβλιογραφία αναδεικνύει προκλήσεις όπως η περιορισμένη διαθεσιμότητα και ποικιλία επισημασμένων δεδομένων, η δυσκολία γενίκευσης των μοντέλων σε διαφορετικά έργα και τεχνολογίες, η ανάγκη για ερμηνευσιμότητα των συστημάτων βαθιάς μάθησης και η ηθική διάσταση της ψηφιακής αποκατάστασης (Basu, et al., 2023; O'Brien, et al., 2023; Kumar & Gupta, 2025). Η ανάγκη για διαφάνεια, επεξηγησιμότητα και συνεργασία ανθρώπου-μηχανής καθίσταται κρίσιμη για την ευρύτερη αποδοχή και εφαρμογή των αυτοματοποιημένων συστημάτων στην πράξη (Fontaine, 2024).

1.4. Στόχοι, ερευνητικά ερωτήματα και συμβολή της παρούσας εργασίας

Η παρούσα πτυχιακή εργασία επικεντρώνεται στην ανάπτυξη και αξιολόγηση μιας εφαρμογής για την ανίχνευση και αξιολόγηση φθορών σε πίνακες ζωγραφικής, αξιοποιώντας τεχνικές μηχανικής μάθησης και σύγχρονες μεθόδους βαθιάς μάθησης. Η

υλοποίηση βασίζεται σε δημόσια διαθέσιμα σύνολα δεδομένων και στοχεύει στην ενίσχυση της τεχνολογικής υποστήριξης της συντήρησης έργων τέχνης.

Οι βασικοί στόχοι της εργασίας είναι:

- Η διερεύνηση και κριτική αποτίμηση της διεθνούς βιβλιογραφίας σχετικά με την αυτόματη ανίχνευση και αξιολόγηση φθορών σε έργα τέχνης,
- Η επιλογή κατάλληλων δεδομένων, εργαλείων και μεθοδολογίας για την υλοποίηση μιας εφαρμογής ανίχνευσης και αξιολόγησης ζημιών,
- Η ανάπτυξη και αξιολόγηση ενός μοντέλου μηχανικής μάθησης για την ανίχνευση και εκτίμηση φθορών σε ψηφιακές εικόνες πινάκων ζωγραφικής.

Για την επίτευξη των παραπάνω στόχων, η εργασία καθοδηγείται από τα ακόλουθα ερευνητικά ερωτήματα:

1. Ποιες τεχνικές μηχανικής μάθησης είναι πιο αποτελεσματικές για την ανίχνευση και αξιολόγηση φθορών σε έργα τέχνης;
2. Ποια χαρακτηριστικά πρέπει να διαθέτουν τα σύνολα δεδομένων και οι μέθοδοι επισημείωσης ώστε να είναι κατάλληλα για την εκπαίδευση και αξιολόγηση τέτοιων μοντέλων;
3. Με ποιους τρόπους μπορεί να εκτιμηθεί αυτόματα η σοβαρότητα και ο τύπος της φθοράς, και ποια είναι η πρακτική αξία αυτής της εκτίμησης για τους επαγγελματίες συντήρησης;
4. Ποια είναι τα πλεονεκτήματα και οι περιορισμοί της αυτοματοποιημένης ανίχνευσης και αξιολόγησης σε σύγκριση με τις παραδοσιακές μεθόδους συντήρησης;

Τα παραπάνω ερωτήματα και προκλήσεις αναλύονται διεξοδικά στο **Κεφάλαιο 2**, το οποίο παρουσιάζει τη σχετική βιβλιογραφία και τις τεχνικές προσεγγίσεις που θεμελιώνουν τη μεθοδολογία της παρούσας εργασίας (**Κεφάλαιο 3**). Στο **Κεφάλαιο 4** παρουσιάζονται τα αποτελέσματα της υλοποίησης και αξιολόγησης του μοντέλου, συνοδευόμενα από συζήτηση και κριτική αποτίμηση των ευρημάτων. Τέλος, στο **Κεφάλαιο 5** συνοψίζονται τα βασικά συμπεράσματα και προτείνονται κατευθύνσεις για μελλοντική έρευνα και εφαρμογή.

Η παρούσα μελέτη επιδιώκει να αναδείξει τα πλεονεκτήματα της αυτοματοποιημένης ανίχνευσης και αξιολόγησης φθορών σε αντικείμενα πολιτιστικής κληρονομιάς, συμβάλλοντας παράλληλα στην κατανόηση των τεχνικών και πρακτικών προκλήσεων που συνοδεύουν την εφαρμογή σχετικών τεχνολογιών. Επιπλέον, διατυπώνει προτάσεις για τη βελτιστοποίηση μελλοντικών εφαρμογών στον τομέα της συντήρησης, με στόχο την ενίσχυση της αποδοτικότητας και της ακρίβειας των παρεμβάσεων. Ιδιαίτερη έμφαση

δίδεται στις ηθικές, κοινωνικές και θεσμικές διαστάσεις της ενσωμάτωσης της τεχνητής νοημοσύνης στη διαχείριση και διατήρηση έργων τέχνης, αναγνωρίζοντας ότι η τεχνολογία δεν λειτουργεί εν κενώ, αλλά εντάσσεται σε ένα ευρύτερο πλαίσιο πολιτισμικής ευθύνης. Η διερεύνηση αυτών των πτυχών ενισχύει τον διεπιστημονικό χαρακτήρα της παρούσας εργασίας και συμβάλλει στην υπεύθυνη, ολιστική και ουσιαστική ενσωμάτωση καινοτόμων τεχνολογιών στην πρακτική της συντήρησης.

ΚΕΦΑΛΑΙΟ 2 Βιβλιογραφική Επισκόπηση

2.1. Εισαγωγή

Η διατήρηση και αποκατάσταση ζωγραφικών έργων τέχνης αποτελεί ένα διαρκώς εξελισσόμενο πεδίο, όπου η τεχνολογική πρόοδος και η διεπιστημονική συνεργασία διαδραματίζουν καθοριστικό ρόλο στη διαφύλαξη της πολιτιστικής κληρονομιάς. Η ακριβής ανίχνευση και αξιολόγηση φθορών, όπως ρωγμές, *lacunae*, μούχλα και δομικές αλλοιώσεις, είναι απαραίτητη για την ορθή συντήρηση των έργων. Παραδοσιακά, η διαδικασία αυτή βασίζεται στην οπτική επιθεώρηση από ειδικούς, η οποία χαρακτηρίζεται από υποκειμενικότητα και περιορισμένη επαναληψιμότητα, ειδικά σε μεγάλες συλλογές ή έργα με σύνθετες τεχνικές (Huang, et al., 2020; Garcia-Moreno, et al., 2024a).

Η ραγδαία εξέλιξη της μηχανικής μάθησης (ML) και της βαθιάς μάθησης (DL) έχει οδηγήσει στην ανάπτυξη αυτοματοποιημένων συστημάτων που βελτιώνουν την ακρίβεια, την αντικειμενικότητα και την αποδοτικότητα της ανίχνευσης και αξιολόγησης φθορών (Sizyakin, et al., 2018; Sizyakin, et al., 2022; Garcia-Moreno, et al., 2024a). Οι τεχνικές αυτές έχουν ήδη αποδείξει τη χρησιμότητά τους στην ανίχνευση ρωγμών, απώλειας χρώματος, μούχλας και άλλων φθορών, συχνά με ακρίβεια που προσεγγίζει ή και ξεπερνά τις παραδοσιακές μεθόδους (Nordin, et al., 2025; Huang, et al., 2020; Sizyakin, et al., 2022).

2.1.α. Επεξήγηση Τεχνικών και Αλγορίθμων

Η εφαρμογή τεχνητής νοημοσύνης στη συντήρηση έργων τέχνης βασίζεται σε ένα ευρύ φάσμα αλγορίθμων και τεχνικών, οι οποίοι επιτελούν διαφορετικούς ρόλους — από την ανίχνευση φθορών έως την αποκατάσταση εικόνας και την πρόβλεψη μελλοντικής αλλοίωσης. Παρακάτω παρουσιάζονται συνοπτικά οι βασικές τεχνολογίες που αναλύονται στη συνέχεια της βιβλιογραφικής ανασκόπησης:

- **Συνελικτικά Νευρωνικά Δίκτυα (CNNs):** Αποτελούν τη βάση για πολλές εφαρμογές στην ανάλυση εικόνας. Μαθαίνουν να εντοπίζουν χαρακτηριστικά όπως άκρα, υφές και μοτίβα, και χρησιμοποιούνται ευρέως για την ανίχνευση ρωγμών, αποχρωματισμών και άλλων φθορών.
- **U-Net:** Εξειδικευμένη αρχιτεκτονική CNN για τμηματοποίηση εικόνας, ιδιαίτερα αποτελεσματική σε εφαρμογές όπου απαιτείται ακρίβεια σε επίπεδο εικονοστοιχείου (pixel-level), όπως η εντόπιση ρωγμών ή *lacunae* σε πίνακες.
- **R-CNN και Mask R-CNN:**
 - Το **R-CNN (Region-based CNN)** εντοπίζει αντικείμενα σε εικόνες εντοπίζοντας πρώτα περιοχές ενδιαφέροντος (regions of interest) και στη συνέχεια εφαρμόζοντας CNN για ταξινόμηση.
 - Το **Mask R-CNN** είναι επέκταση του Faster R-CNN που προσθέτει δυνατότητα τμηματοποίησης σε επίπεδο pixel, καθιστώντας το ιδανικό για την αναγνώριση και κατηγοριοποίηση φθορών με ακρίβεια.

- **Generative Adversarial Networks (GANs):** Αποτελούνται από δύο δίκτυα – έναν γεννήτορα και έναν διακριτή – που ανταγωνίζονται μεταξύ τους. Χρησιμοποιούνται για την αποκατάσταση κατεστραμμένων περιοχών (inpainting), δημιουργώντας ρεαλιστικές εικόνες που συμπληρώνουν τα κενά.
- **Diffusion Models:** Νεότερη κατηγορία γενετικών μοντέλων που δημιουργούν εικόνες μέσω σταδιακής αφαίρεσης θορύβου. Παρέχουν υψηλής ποιότητας αποτελέσματα, αλλά απαιτούν μεγάλα datasets και σημαντική υπολογιστική ισχύ.
- **Sparse Coding:** Τεχνική αναπαράστασης εικόνας ως γραμμικός συνδυασμός λίγων βασικών στοιχείων (atoms). Χρησιμοποιείται για την ανίχνευση απώλειας χρώματος, ειδικά σε πολυτροπικά δεδομένα (π.χ. UV, IR).
- **XGBoost:** Αλγόριθμος ενισχυτικής μάθησης (gradient boosting) που χρησιμοποιείται για προβλεπτικά μοντέλα, όπως η εκτίμηση της εξέλιξης ρωγμών ή άλλων φθορών με βάση ιστορικά και φυσικά δεδομένα.
- **Active Learning:** Στρατηγική επιλογής των πιο πληροφοριακών δειγμάτων για επισημείωση (annotation), μειώνοντας το κόστος labeling χωρίς σημαντική απώλεια απόδοσης. Ιδιαίτερα χρήσιμη σε περιβάλλοντα με περιορισμένα δεδομένα.
- **Finite Element Analysis (FEA):** Μηχανική μέθοδος προσομοίωσης που επιτρέπει την ανάλυση μηχανικών καταπονήσεων και την πρόβλεψη φθοράς σε υλικά και δομές. Όταν συνδυάζεται με ML, επιτρέπει την προληπτική συντήρηση.

Η κατανόηση αυτών των τεχνικών είναι απαραίτητη για την ερμηνεία των αποτελεσμάτων που παρουσιάζονται στις επόμενες ενότητες, καθώς και για την αξιολόγηση της εφαρμοσιμότητάς τους στην πράξη.

2.2. Κριτική Αποτίμηση Τεχνικών και Αποτελεσμάτων

2.2.α. Ανίχνευση Φθορών με Πολυτροπικά Δεδομένα και CNNs

Η ανίχνευση φθορών σε ζωγραφικά έργα τέχνης έχει επωφεληθεί σημαντικά από την αξιοποίηση πολυτροπικών δεδομένων και την εφαρμογή συνελκτικών νευρωνικών δικτύων (CNNs). Η εργασία των Huang, et al. (2020) εισήγαγε μια καινοτόμο προσέγγιση που συνδυάζει εικόνες από διαφορετικά φάσματα (ορατό, υπέρυθρο, υπεριώδες) με τεχνικές sparse coding, επιτρέποντας την ανίχνευση απώλειας χρώματος και υποκείμενων φθορών που δεν είναι ορατές στο φάσμα του ορατού φωτός. Η πολυτροπική ανάλυση ενισχύει την πληρότητα της διάγνωσης, αλλά απαιτεί εξειδικευμένο εξοπλισμό και υψηλό επίπεδο τεχνικής κατάρτισης, περιορίζοντας την ευρεία εφαρμογή της.

Παράλληλα οι Basu, et al. (2023) υπογραμμίζουν τη σημασία της διαλειτουργικότητας μεταξύ διαφορετικών τεχνικών απεικόνισης (π.χ. RTI, XRF, IRR), επισημαίνοντας ότι η επιτυχία των data-driven προσεγγίσεων εξαρτάται από την ικανότητα ενοποίησης ετερογενών δεδομένων. Ωστόσο, η πολυπλοκότητα αυτής της ενοποίησης και η ανάγκη για υψηλής ποιότητας επισημείωση (annotation) αποτελούν σημαντικές προκλήσεις.

Η εφαρμογή συνελκτικών νευρωνικών δικτύων (CNNs), και ειδικότερα του **U-Net**, έχει αποδειχθεί ιδιαίτερα αποτελεσματική στην ανίχνευση ρωγμών και άλλων μορφών φθοράς. Οι Sizyakin et al. (2018) και Dulecha, et al. (2019) αξιολόγησαν την απόδοση των CNNs σε εικόνες με διαφορετικές γωνίες φωτισμού, χρησιμοποιώντας τεχνικές RTI (Reflectance Transformation Imaging). Τα αποτελέσματα έδειξαν ακρίβεια σε επίπεδο εικονοστοιχείου (pixel-level accuracy) άνω του 92%, επιβεβαιώνοντας την ικανότητα των CNNs να εντοπίζουν λεπτές ρωγμές ακόμη και σε επιφάνειες με έντονη υφή.

Η ανθεκτικότητα (robustness) των μοντέλων αυτών δοκιμάστηκε περαιτέρω από τους Sandoval, et al. (2020), οι οποίοι χρησιμοποίησαν εικόνες με τεχνητές ή προσομοιωμένες φθορές. Τα CNNs διατήρησαν υψηλή απόδοση ακόμη και σε περιπτώσεις με σημαντικές απώλειες πληροφορίας, αν και η ακρίβεια μειώθηκε όταν οι φθορές απέκρυπταν βασικά μορφολογικά χαρακτηριστικά του έργου.

Συνολικά, η συνδυαστική χρήση πολυτροπικών δεδομένων και CNNs προσφέρει μια ισχυρή προσέγγιση για την ανίχνευση φθορών, με δυνατότητες γενίκευσης σε διαφορετικά έργα και τεχνολογίες. Ωστόσο, η επιτυχία αυτών των μεθόδων εξαρτάται από την ποιότητα των δεδομένων, την τεχνική υποδομή και την ερμηνευσιμότητα των αποτελεσμάτων, στοιχεία που θα αναλυθούν περαιτέρω στην ενότητα 2.3

2.2.β. Αποκατάσταση Εικόνας και Ανοιχτά Εργαλεία

Η αποκατάσταση κατεστραμμένων περιοχών σε έργα τέχνης μέσω τεχνικών inpainting αποτελεί ένα από τα πιο εντυπωσιακά πεδία εφαρμογής της τεχνητής νοημοσύνης στη συντήρηση. Οι Sizyakin, et al. (2022) και οι Kumar & Gupta (2025) εξετάζουν τη χρήση Generative Adversarial Networks (GANs) και autoencoders για την ανακατασκευή περιοχών με απώλεια υλικού ή χρώματος. Τα GANs, μέσω της αλληλεπίδρασης γεννήτορα και διακριτή, επιτυγχάνουν ρεαλιστικά αποτελέσματα, αλλά συχνά παράγουν "υπερβολικά τέλεια" εικόνες που ενδέχεται να αποκλίνουν από το αρχικό ύφος του έργου.

Τα diffusion models, μια νεότερη γενιά γενετικών μοντέλων, προσφέρουν ακόμη πιο ρεαλιστική ανακατασκευή, ξεκινώντας από θόρυβο και σταδιακά δημιουργώντας εικόνες μέσω αντιστροφής της διαδικασίας διάχυσης. Παρότι υπερέρχουν σε ποιότητα, απαιτούν μεγάλα και ποικίλα datasets και σημαντική υπολογιστική ισχύ, γεγονός που περιορίζει την πρακτική τους εφαρμογή σε μικρότερες συλλογές ή ιδρύματα με περιορισμένους πόρους.

Η εργασία των O'Brien, et al. (2023), που επικεντρώνεται στην ψηφιακή ανακατασκευή του έργου του Antoine François Callet, αναδεικνύει τα όρια και τις ηθικές προεκτάσεις των τεχνικών αυτών. Παρότι τα αποτελέσματα είναι εντυπωσιακά, η χρήση τους σε ιστορικά έργα απαιτεί αυξημένη προσοχή, ώστε να μην αλλοιωθεί η αυθεντικότητα και η ιστορική αξία του έργου.

Η αξιολόγηση της απόδοσης των παραπάνω τεχνικών εξαρτάται σε μεγάλο βαθμό από την ύπαρξη δημόσιων, επισημειωμένων συνόλων δεδομένων. Το ARTeFACT Dataset (Ivanova, et al., 2025) αποτελεί ένα από τα πιο πλήρη και εξειδικευμένα datasets για την ανίχνευση και κατηγοριοποίηση φθορών σε πίνακες ζωγραφικής. Περιλαμβάνει πολυτροπικά δεδομένα, επισημειωμένες περιοχές φθοράς (π.χ. lacunae, stucco, ρωγμές) και χρησιμοποιείται ευρέως για την εκπαίδευση και αξιολόγηση μοντέλων segmentation και inpainting. Η ύπαρξη τέτοιων συνόλων δεδομένων είναι κρίσιμη για την αντικειμενική σύγκριση διαφορετικών προσεγγίσεων και την επαναληψιμότητα των αποτελεσμάτων.

Στο πλαίσιο της πρακτικής εφαρμογής, το ανοιχτό λογισμικό Nirvan101/Image-Restoration-deep-learning (2025) προσφέρει ένα λειτουργικό περιβάλλον για την αξιολόγηση διαφορετικών τεχνικών αποκατάστασης, όπως GANs και diffusion models. Η συλλογή περιλαμβάνει 68 εικόνες, εκ των οποίων οι 20 είναι σε καλή κατάσταση, ενώ οι υπόλοιπες παρουσιάζουν φθορές όπως αποχρωματισμένα σημεία και απώλεια υλικού. Το λογισμικό αυτό επιτρέπει τη συγκριτική αξιολόγηση διαφορετικών μεθόδων και μπορεί να χρησιμοποιηθεί τόσο για ερευνητικούς σκοπούς όσο και για εκπαιδευτική χρήση σε προγράμματα συντήρησης.

Η ενσωμάτωση τέτοιων εργαλείων στην καθημερινή πρακτική των συντηρητών εξαρτάται από τη διαφάνεια των αποτελεσμάτων, τη δυνατότητα ελέγχου των παραμέτρων και την ερμηνευσιμότητα των προτάσεων αποκατάστασης. Αυτά τα ζητήματα θα αναλυθούν περαιτέρω στην ενότητα 2.3.

2.2.γ. Επισημείωση, Εξειδίκευση και Επέκταση Δεδομένων

Η ποιότητα και η ποικιλία των δεδομένων εκπαίδευσης αποτελούν κρίσιμους παράγοντες για την επιτυχία των μοντέλων μηχανικής μάθησης στη συντήρηση έργων τέχνης. Η δημιουργία μεγάλων και αξιόπιστων συνόλων δεδομένων είναι ιδιαίτερα απαιτητική, καθώς απαιτεί εξειδικευμένη γνώση, χρόνο και ανθρώπινους πόρους για την επισημείωση (annotation) των φθορών.

Η προσέγγιση Deep Active Learning (DAL), όπως παρουσιάστηκε από τους Nadisic, et al. (2024) στο πλαίσιο του συστήματος DAL4ART, προσφέρει μια λύση στο πρόβλημα αυτό. Μέσω της ενεργής μάθησης, το σύστημα επιλέγει αυτόματα τα πιο πληροφοριακά δείγματα για επισημείωση, μειώνοντας το κόστος labeling έως και 70% χωρίς σημαντική απώλεια απόδοσης. Η τεχνική αυτή είναι ιδιαίτερα χρήσιμη για την επέκταση των datasets σε νέα έργα, σπάνιες φθορές ή υποεκπροσωπούμενες τεχνοτροπίες.

Ωστόσο, η ποιότητα των αρχικών labels παραμένει καθοριστική. Όπως επισημαίνουν οι Su, et al. (2022), τα deep learning μοντέλα είναι ιδιαίτερα ευαίσθητα σε σφάλματα επισημείωσης, τα οποία μπορούν να οδηγήσουν σε συστηματικές αποκλίσεις στα

αποτελέσματα. Η ενεργή μάθηση μπορεί να μειώσει το κόστος, αλλά δεν υποκαθιστά την ανάγκη για συνεργασία με ειδικούς συντηρητές κατά τη φάση της επισημείωσης.

Η ανάγκη για εξειδίκευση ανά υλικό και τύπο φθοράς αναδεικνύεται και από τη μελέτη των Nordin, et al. (2025), οι οποίοι συνέκριναν μοντέλα μηχανικής μάθησης (CART, LDA) με rule-based συστήματα για την ανίχνευση και κατηγοριοποίηση μυκητιακής ανάπτυξης (μούχλας) σε πίνακες ζωγραφικής. Τα ML μοντέλα υπερείχαν σε ακρίβεια και προσαρμοστικότητα, αλλά η επιτυχία τους εξαρτήθηκε από την ποιότητα και την εξειδίκευση των χαρακτηριστικών που χρησιμοποιήθηκαν για την εκπαίδευση.

Παρόμοια, οι Müller et al. (2021) ανέδειξαν τη σημασία της επιλογής κατάλληλων χαρακτηριστικών υφής και ταξινομητών (π.χ. SVM, LDA) ανάλογα με το υλικό του έργου (π.χ. ξύλο, καμβάς, τοιχογραφία). Η προσαρμογή των μοντέλων στις ιδιαιτερότητες κάθε έργου είναι απαραίτητη για την επίτευξη υψηλής ακρίβειας και αξιοπιστίας.

Συνολικά, η επέκταση και εξειδίκευση των δεδομένων αποτελεί θεμέλιο για την επιτυχία των συστημάτων τεχνητής νοημοσύνης στη συντήρηση τέχνης. Η ενεργή μάθηση, σε συνδυασμό με τη συνεργασία ειδικών και την προσεκτική επιλογή χαρακτηριστικών, μπορεί να μειώσει το κόστος και να ενισχύσει τη γενίκευση των μοντέλων, προετοιμάζοντας το έδαφος για πιο ευρεία και υπεύθυνη εφαρμογή.

2.2.δ. Ενσωματωμένα Εργαλεία και Προληπτική Συντήρηση

Η μετάβαση από πειραματικά μοντέλα σε πλήρως ενσωματωμένα εργαλεία αποτελεί βασικό βήμα για την πρακτική εφαρμογή της τεχνητής νοημοσύνης στη συντήρηση έργων τέχνης. Το ARTDET των Garcia-Moreno, et al. (2024a) αποτελεί χαρακτηριστικό παράδειγμα ενός τέτοιου εργαλείου, το οποίο ενσωματώνει το Mask R-CNN για την ανίχνευση και κατηγοριοποίηση φθορών όπως lacunae και stucco. Το σύστημα προσφέρει αυτόματη εκτίμηση σοβαρότητας, φιλική διεπαφή χρήστη και δυνατότητα εφαρμογής σε πραγματικές συνθήκες εργασίας, καθιστώντας το ιδιαίτερα χρήσιμο για συντηρητές χωρίς προηγούμενη εμπειρία σε τεχνολογίες AI.

Η επιτυχία τέτοιων εργαλείων εξαρτάται σε μεγάλο βαθμό από τη διαφάνεια των αποτελεσμάτων και τη δυνατότητα ερμηνείας των προβλέψεων. Όπως επισημαίνουν οι Basu et al. (2023) και David Stork (2023), η αποδοχή από την κοινότητα των συντηρητών προϋποθέτει ότι τα συστήματα τεχνητής νοημοσύνης λειτουργούν ως υποστηρικτικά εργαλεία και όχι ως αυτόνομοι αντικαταστάτες της ανθρώπινης κρίσης.

Παράλληλα, η εργασία των Califano, et al. (2022) εισάγει μια υβριδική προσέγγιση που συνδυάζει μοντέλα μηχανικής μάθησης (XGBoost) με μηχανική προσομοίωση (finite element analysis) για την πρόβλεψη της εξέλιξης φθοράς, όπως η διάδοση ρωγμών. Η δυνατότητα

πρόβλεψης επιτρέπει την προληπτική συντήρηση, μειώνοντας την ανάγκη για επεμβατικές παρεμβάσεις και βελτιώνοντας τη διαχείριση των συλλογών.

Η εφαρμογή τέτοιων μοντέλων απαιτεί εξειδικευμένα δεδομένα ανά έργο και ακριβή αρχική καταγραφή των φυσικών ιδιοτήτων του υλικού. Ωστόσο, η ενσωμάτωση φυσικών και υπολογιστικών μοντέλων ανοίγει νέους δρόμους για τη συστηματική παρακολούθηση και μακροπρόθεσμη διατήρηση των έργων τέχνης.

Συνολικά, η ανάπτυξη εργαλείων που συνδυάζουν τεχνική ακρίβεια, χρησιμότητα και ερμηνευσιμότητα αποτελεί βασική προϋπόθεση για την επιτυχή ενσωμάτωση της τεχνητής νοημοσύνης στην καθημερινή πρακτική της συντήρησης.

2.2.ε. Ηθικές και Διεπιστημονικές Προεκτάσεις

Η ενσωμάτωση τεχνητής νοημοσύνης στη συντήρηση έργων τέχνης δεν είναι μόνο τεχνικό ζήτημα, αλλά εγείρει και σημαντικά ηθικά, ερμηνευτικά και κοινωνικά ερωτήματα. Οι O'Brien, et al. (2023) και Stork (2023) επισημαίνουν ότι η χρήση γενετικών μοντέλων, όπως τα GANs και diffusion models, μπορεί να οδηγήσει σε ψηφιακές ανακατασκευές που υπερβαίνουν το αρχικό έργο, αλλοιώνοντας την καλλιτεχνική πρόθεση και την ιστορική αυθεντικότητα. Η τεχνολογία, αν και ισχυρή, δεν μπορεί να αντικαταστήσει την ανθρώπινη κρίση σε ζητήματα αισθητικής, πολιτισμικής σημασίας και ιστορικής ερμηνείας.

Η ανάγκη για διαφάνεια και ερμηνευσιμότητα των αποτελεσμάτων είναι κρίσιμη για την αποδοχή των συστημάτων AI από την κοινότητα των συντηρητών. Οι Fontaine (2024), Adadi & Berrada (2018) και Molina & Sundar (2022) τονίζουν ότι η Explainable AI (XAI) είναι απαραίτητη για να κατανοούν οι χρήστες πώς και γιατί ένα μοντέλο κατέληξε σε ένα συγκεκριμένο αποτέλεσμα. Η έλλειψη εξήγησης μπορεί να οδηγήσει σε δυσπιστία, απόρριψη ή λανθασμένη χρήση των εργαλείων. (Adadi & Berrada, 2018; Molina & Sundar, 2022; Mahalle & Ingle, 2024).

Επιπλέον, η διεπιστημονική συνεργασία μεταξύ συντηρητών, επιστημόνων υπολογιστών, ιστορικών τέχνης και ηθικολόγων είναι απαραίτητη για την υπεύθυνη ανάπτυξη και εφαρμογή αυτών των τεχνολογιών. Η εργασία των Puy-Alquiza, et al. (2021) δείχνει πώς τεχνικές mapping και ανάλυσης φθορών από τον τομέα της αρχιτεκτονικής συντήρησης μπορούν να μεταφερθούν και να προσαρμοστούν στη ζωγραφική, ενισχύοντας τη μεταφορά τεχνογνωσίας και την καινοτομία.

Συνολικά, η επιτυχής ενσωμάτωση της τεχνητής νοημοσύνης στη συντήρηση τέχνης απαιτεί όχι μόνο τεχνική αρτιότητα, αλλά και σεβασμό στην πολιτιστική αξία, διαφάνεια στη διαδικασία και συνεχή διάλογο μεταξύ των εμπλεκόμενων επιστημονικών και επαγγελματικών κοινοτήτων.

2.3. Ερμηνεία και Εφαρμοσιμότητα των Ευρημάτων

Η ανάλυση της βιβλιογραφίας αποκαλύπτει ότι οι τεχνικές τεχνητής νοημοσύνης προσφέρουν σημαντικές δυνατότητες για την ενίσχυση της ακρίβειας, της αποδοτικότητας και της τεκμηρίωσης στη συντήρηση ζωγραφικών έργων τέχνης. Ωστόσο, η πρακτική εφαρμογή τους εξαρτάται από μια σειρά παραγόντων που σχετίζονται με την ερμηνευσιμότητα, τη γενίκευση, την αποδοχή από τους επαγγελματίες και την ενσωμάτωσή τους σε πραγματικά περιβάλλοντα εργασίας.

2.3.α. Πρακτική Αξία και Υποστήριξη Απόφασης

Τα αυτοματοποιημένα εργαλεία, όπως το ARTDET (Garcia-Moreno, et al., 2024a), προσφέρουν δυνατότητες προ-ανάλυσης, ταξινόμησης σοβαρότητας φθοράς και εικονικής αποκατάστασης, επιτρέποντας στους συντηρητές να πειραματιστούν με εναλλακτικά σενάρια πριν από την υλική παρέμβαση. Η χρήση εργαλείων όπως το Nirvan101/Image-Restoration-deep-learning και datasets όπως το ARTeFACT ενισχύει τη δυνατότητα εφαρμογής των μοντέλων σε πραγματικά έργα, προσφέροντας ένα πλαίσιο για τεκμηριωμένη λήψη αποφάσεων.

2.3.β. Ερμηνευσιμότητα και Εμπιστοσύνη

Η ερμηνευσιμότητα (explainability) παραμένει κρίσιμος παράγοντας για την αποδοχή των συστημάτων από την κοινότητα των συντηρητών. Εργαλεία όπως τα saliency maps, Grad-CAM και attention mechanisms επιτρέπουν την οπτικοποίηση των περιοχών που επηρεάζουν τις αποφάσεις του μοντέλου (Basu, et al., 2023; Sizyakin, et al., 2022). Η δυνατότητα εξήγησης ενισχύει την εμπιστοσύνη και διευκολύνει τη συνεργασία ανθρώπου-μηχανής, ειδικά σε περιβάλλοντα όπου η τελική απόφαση πρέπει να παραμένει στον ειδικό (Stork, 2023).

2.3.γ. Γενίκευση και Ανθεκτικότητα Μοντέλων

Η γενίκευση των μοντέλων σε διαφορετικά έργα, τεχνοτροπίες και τύπους φθοράς αποτελεί συνεχή πρόκληση. Η περιορισμένη ποικιλία στα διαθέσιμα datasets, όπως επισημαίνουν οι Garcia-Moreno et al. (2024b) και Nordin et al. (2025), επηρεάζει την απόδοση των μοντέλων σε νέα δεδομένα. Τεχνικές όπως το transfer learning, η domain adaptation και η semi-supervised learning προτείνονται για τη βελτίωση της γενικευσιμότητας (Su, et al., 2022; Basu, et al., 2023). Επιπλέον, η χρήση active learning (Nadisic, et al., 2024) μπορεί να μειώσει το κόστος επισημείωσης και να επιταχύνει την επέκταση των συνόλων δεδομένων.

2.3.δ Συγκριτική Ανάλυση Επιλεγμένων Μελετών

Η ενότητα αυτή παρουσιάζει μια συγκριτική ανάλυση έξι σημαντικών μελετών που αφορούν την εφαρμογή τεχνικών τεχνητής νοημοσύνης στη συντήρηση έργων τέχνης. Ο

πίνακας που ακολουθεί συνοψίζει τις τεχνικές λεπτομέρειες, τα δεδομένα που χρησιμοποιήθηκαν, τα αποτελέσματα και τα βασικά συμπεράσματα κάθε μελέτης.

Τεχνικές Περιλήψεις

[Sizyakin et al. \(2018\)](#)

Η μελέτη αυτή εφάρμοσε U-Net convolutional neural networks για την ανίχνευση ρωγμών σε πίνακες ζωγραφικής με χρήση Reflectance Transformation Imaging (RTI). Το μοντέλο πέτυχε ακρίβεια σε επίπεδο pixel πάνω από 92%, αποδεικνύοντας την αποτελεσματικότητα της βαθιάς μάθησης στην υψηλής ανάλυσης εντοπισμό ζημιών.

[Garcia-Moreno et al. \(2024a\)](#)

Το έργο ARTDET εισήγαγε ένα σύστημα βασισμένο στο Mask R-CNN για την ανίχνευση lacunae και άλλων τύπων ζημιών σε πίνακες ζωγραφικής. Χρησιμοποιώντας ένα προσαρμοσμένο σύνολο δεδομένων, το μοντέλο πέτυχε μέση ακρίβεια (mAP) περίπου 0.78, προσφέροντας ένα ισχυρό ανοιχτό εργαλείο για τους συντηρητές.

[Huang et al. \(2020\)](#)

Η εργασία αυτή συνδύασε sparse coding με CNNs για την ανίχνευση απώλειας χρώματος σε πολυτροπικές εικόνες (IR, UV, ορατές). Το μοντέλο πέτυχε F1-score 0.85, υπογραμμίζοντας τα οφέλη της συγχώνευσης πολυτροπικών δεδομένων στη βελτίωση της ακρίβειας ανίχνευσης.

[Kumar & Gupta \(2025\)](#)

Μια συγκριτική μελέτη των GANs και diffusion models για την εικονική αποκατάσταση κατεστραμμένων έργων τέχνης. Τα diffusion models υπερτερούν των GANs σε δομική ομοιότητα (SSIM > 0.90), προσφέροντας πιο ρεαλιστικές και χωρίς τεχνουργήματα αποκαταστάσεις.

[Nordin et al. \(2025\)](#)

Η εργασία αυτή συνέκρινε rule-based συστήματα με XGBoost classifiers για την ανίχνευση μούχλας σε πίνακες ζωγραφικής. Το ML μοντέλο πέτυχε ακρίβεια 94%, υπερτερώντας σημαντικά της rule-based προσέγγισης (78%), τονίζοντας την αξία των μεθόδων που βασίζονται σε δεδομένα.

[O'Brien et al. \(2023\)](#)

Εξετάζει τη χρήση diffusion models για την ψηφιακή αποκατάσταση ιστορικών πινάκων ζωγραφικής. Παρόλο που η αξιολόγηση ήταν ποιοτική, η μελέτη έθεσε σημαντικά ηθικά και ερμηνευτικά ερωτήματα σχετικά με τον ρόλο της AI στην αποκατάσταση τέχνης.

Η μελέτη αυτή εισάγει το ARTeFACT, ένα πολυδιάστατο benchmark dataset που περιλαμβάνει 15 τύπους φθοράς σε έργα τέχνης διαφορετικών υλικών, όπως πίνακες, τοιχογραφίες και φωτογραφίες. Οι συγγραφείς αξιολόγησαν CNNs, Transformers και diffusion models σε σενάρια cross-media, επιτυγχάνοντας mIoU έως 0.81. Το ARTeFACT ξεχωρίζει για την ποικιλομορφία και την ποιότητα των επισημειώσεων, προσφέροντας ένα ρεαλιστικό και απαιτητικό πλαίσιο για τη γενίκευση και τη συγκριτική αξιολόγηση μοντέλων.

Κριτική Σύγκριση

Οι ανασκοπούμενες μελέτες συλλογικά καταδεικνύουν την αυξανόμενη πολυπλοκότητα των τεχνικών AI στην συντήρηση τέχνης. Τα μοντέλα βασισμένα σε CNN, όπως το U-Net και το Mask R-CNN, έχουν αποδειχθεί ιδιαίτερα αποτελεσματικά για εργασίες τμηματοποίησης, ιδιαίτερα στην ανίχνευση ρωγμών και lacunae. Οι πολυτροπικές προσεγγίσεις (Huang, et al., 2020) ενισχύουν την ακρίβεια ανίχνευσης μέσω της συγχώνευσης διαφορετικών απεικονιστικών μεθόδων. Εν τω μεταξύ, τα γεννητικά μοντέλα (Kumar & Gupta, 2025; O'Brien, et al., 2023) δείχνουν υποσχόμενα αποτελέσματα στην εικονική αποκατάσταση, αν και εγείρουν ερμηνευτικά ζητήματα. Η σύγκριση μεταξύ rule-based και ML συστημάτων (Nordin, et al., 2025) υπογραμμίζει την ανωτερότητα των μεθόδων που βασίζονται σε δεδομένα.

Πίνακας 1. Τεχνική Σύγκριση Βασικών Ακαδημαϊκών Εργασιών για την Τεχνητή Νοημοσύνη στη Συντήρηση Έργων Τέχνης

Παραπομπή	Τεχνική(ές)	Εφαρμογή / Σύνολο Δεδομένων	Αποτελέσματα Απόδοσης	Κύρια Συμπεράσματα
Sizyakin et al., 2018	CNNs (U-Net)	Ανίχνευση ρωγμών σε πίνακες με χρήση RTI εικόνων	Ακρίβεια σε επίπεδο pixel >92%	Πρώιμη εφαρμογή του U-Net στην πολιτιστική κληρονομιά; ισχυρή απόδοση σε δεδομένα RTI
Garcia-Moreno et al., 2024 (ARTDET)	Mask R-CNN	Πίνακες ζωγραφικής (ARTDET σύνολο δεδομένων)	mAP $\approx$ 0.78 για ανίχνευση lacunae	Ισχυρή τμηματοποίηση τύπων ζημιών; ανοιχτό εργαλείο
Huang et al., 2020	Sparse coding + CNNs	Ανίχνευση απώλειας χρώματος σε πολυτροπικές εικόνες	F1-score $\approx$ 0.85	Συνδυασμός δεδομένων IR/UV/ορατών; ισχυρή συγχώνευση πολυτροπικών δεδομένων
Kumar & Gupta, 2025	GANs, Diffusion Models	Εικονική αποκατάσταση κατεστραμμένων έργων τέχνης	SSIM > 0.90	Υψηλής πιστότητας αποκατάσταση; τα diffusion models υπερτερούν των GANs σε ρεαλισμό
Nordin et al., 2025	Rule-based vs ML (XGBoost)	Ταξινόμηση μούχλας σε πίνακες ζωγραφικής	Ακρίβεια ML $\approx$ 94% vs Rule-based $\approx$ 78%	Τα ML μοντέλα υπερτερούν σημαντικά των rule-based συστημάτων
O'Brien et al., 2023	Diffusion models	Ψηφιακή αποκατάσταση ιστορικών πινάκων ζωγραφικής	Ποιοτική αξιολόγηση	Εξετάζει τα ηθικά και ερμηνευτικά όρια της AI στην αποκατάσταση
Ivanova et al., 2025	CNNs, Transformers, Diffusion Models	ARTeFACT: 15 τύποι φθοράς σε πολλαπλά υλικά	mIoU έως 0.81 σε cross-media αξιολόγηση	Πρώτο benchmark dataset για γενίκευση σε διαφορετικά αναλογικά μέσα

Τέλος, η συμβολή της Ivanova et al. (2025) με το ARTeFACT dataset αναδεικνύει τη σημασία της γενίκευσης και της αξιολόγησης σε ετερογενή δεδομένα, προσφέροντας ένα νέο πρότυπο για benchmarking στην πολιτιστική τεχνητή νοημοσύνη. Συνολικά, αυτές οι εργασίες αναδεικνύουν μια στροφή προς εξηγήσιμες, βασισμένες σε δεδομένα και ηθικά ευαισθητοποιημένες εφαρμογές AI στην πολιτιστική κληρονομιά.

2.3.ε. Ενσωμάτωση σε Ροές Εργασίας και Ανθρώπινη Συνεργασία

Η επιτυχής εφαρμογή των συστημάτων AI εξαρτάται από την ενσωμάτωσή τους στις υπάρχουσες ροές εργασίας των συντηρητών. Η προσέγγιση human-in-the-loop, όπου ο ειδικός έχει τον τελικό λόγο, θεωρείται η βέλτιστη πρακτική για τη διασφάλιση της ποιότητας και της ηθικής της αποκατάστασης (Stork, 2023). Η διεπιστημονική συνεργασία μεταξύ συντηρητών, επιστημόνων υπολογιστών και ιστορικών τέχνης είναι απαραίτητη για την προσαρμογή των εργαλείων στις ανάγκες της πράξης.

Παρά τις σημαντικές δυνατότητες που προσφέρουν τα συστήματα τεχνητής νοημοσύνης στη συντήρηση έργων τέχνης, η εφαρμογή τους συνοδεύεται από θεωρητικές, τεχνικές και μεθοδολογικές προκλήσεις. Η κατανόηση αυτών των περιορισμών είναι απαραίτητη για την υπεύθυνη και αποτελεσματική αξιοποίησή τους στην πράξη.

2.4. Θεωρητικές, Μεθοδολογικές και Τεχνικές Προκλήσεις

Η εφαρμογή τεχνικών μηχανικής μάθησης στη συντήρηση έργων τέχνης συνοδεύεται από σημαντικές προκλήσεις που αφορούν τόσο τη θεωρητική τεκμηρίωση όσο και την πρακτική υλοποίηση των συστημάτων.

2.4.α. Περιορισμοί των Υφιστάμενων Μεθόδων

Παρά την πρόοδο, οι υφιστάμενες μέθοδοι παρουσιάζουν κρίσιμους περιορισμούς:

- **Έλλειψη ποικιλίας και ποσότητας δεδομένων:** Τα διαθέσιμα datasets είναι συχνά περιορισμένα σε μέγεθος, μη ισορροπημένα ή μη επαρκώς επισημασμένα, γεγονός που επηρεάζει τη γενικευσιμότητα των μοντέλων (Basu, et al., 2023). Η έλλειψη μεγάλων, ποικίλων (και ισορροπημένων) και επισημειωμένων datasets περιορίζει τη γενικευσιμότητα των μοντέλων (Garcia-Moreno, et al., 2024b; Ivanova, et al., 2025).
- **Ανάγκη για εξειδικευμένο εξοπλισμό:** Η συλλογή δεδομένων υψηλής ποιότητας (π.χ. υπέρυθρη φασματοσκοπία, X-ray fluorescence) απαιτεί ακριβό και δύσκολα προσβάσιμο εξοπλισμό, περιορίζοντας την ευρεία εφαρμογή των μεθόδων. Η χρήση πολυτροπικών δεδομένων (UV, IR, XRF) απαιτεί τεχνική υποδομή που δεν είναι διαθέσιμη σε όλα τα ιδρύματα (Huang, et al., 2020).
- **Black-box φύση των DL μοντέλων:** Η έλλειψη διαφάνειας στις αποφάσεις των μοντέλων δημιουργεί δυσπιστία στους επαγγελματίες συντηρητές και δυσκολεύει την αποδοχή τους στην πράξη (Adadi & Berrada, 2018; Mahalle & Ingle, 2024).

- **Ανάγκη για explainable AI (XAI):** Η ερμηνευσιμότητα είναι κρίσιμη για την κατανόηση και αποδοχή των αποτελεσμάτων από τους ειδικούς, ιδιαίτερα σε περιβάλλοντα όπου η τελική απόφαση έχει πολιτιστική και ηθική βαρύτητα (O'Brien, et al., 2023).
- **Ευαισθησία σε σφάλματα επισημείωσης:** Τα DL μοντέλα είναι ευάλωτα σε εσφαλμένα labels, γεγονός που μπορεί να οδηγήσει σε συστηματικά σφάλματα (Su, et al., 2022).

2.4.β. Μεθοδολογικές Προτάσεις για Μελλοντική Έρευνα

Για την υπέρβαση των παραπάνω περιορισμών, η βιβλιογραφία προτείνει τις εξής κατευθύνσεις:

- **Ενσωμάτωση πολυτροπικών δεδομένων (multimodal data):** Ο συνδυασμός εικόνων, φασματικών δεδομένων και ιστορικών πληροφοριών μπορεί να ενισχύσει την ακρίβεια και την αξιοπιστία των μοντέλων (Califano, et al., 2022).
- **Ανάπτυξη μεγαλύτερων και πιο ποικίλων datasets:** Η δημιουργία διεθνών, ανοιχτών βάσεων δεδομένων με επισημασμένα παραδείγματα φθοράς αποτελεί προτεραιότητα για την πρόοδο του πεδίου (Puy-Alquiza, et al., 2021).
- **Explainable AI και human-in-the-loop:** Η ανάπτυξη εργαλείων που επιτρέπουν την ερμηνεία των αποτελεσμάτων και τη συνεργασία με τον ειδικό συντηρητή είναι κρίσιμη για την πρακτική εφαρμογή (Basu, et al., 2023).
- **Συνδυασμός ML με φυσικά/μηχανικά μοντέλα:** Η ενσωμάτωση φυσικών μοντέλων (όπως finite element analysis) μπορεί να ενισχύσει την προβλεπτική ικανότητα των συστημάτων και να προσφέρει φυσικά ερμηνεύσιμα αποτελέσματα (Basu, et al., 2023; Califano, et al., 2022).
- **Διεπιστημονική συνεργασία:** Η επιτυχής εφαρμογή απαιτεί τη συνεργασία ειδικών από τομείς όπως η συντήρηση, η πληροφορική, η φυσική και η ιστορία της τέχνης (Ivanova, et al., 2025).
- **Ενεργή μάθηση και semi-supervised learning:** Μείωση κόστους επισημείωσης και επέκταση datasets (Nadisic, et al., 2024).

2.5. Κοινωνικές, Ηθικές και Θεσμικές Πτυχές

Η χρήση τεχνητής νοημοσύνης στη συντήρηση τέχνης εγείρει σημαντικά ζητήματα:

- **Αυθεντικότητα και ιστορική ακεραιότητα:** Τα GANs και diffusion models ενδέχεται να δημιουργήσουν «υπερβολικά τέλειες» αποκαταστάσεις που αλλοιώνουν την πρόθεση του καλλιτέχνη (O'Brien, et al., 2023).
- **Ευθύνη και λήψη αποφάσεων:** Η τελική απόφαση πρέπει να παραμένει στον ειδικό συντηρητή, εντός πλαισίου human-in-the-loop (Stork, 2023).
- **Διαφάνεια και αποδοχή:** Η ερμηνευσιμότητα των αποτελεσμάτων είναι κρίσιμη για την αποδοχή από την κοινότητα (Fontaine, 2024; Molina & Sundar, 2022).

- **Θεσμικά και νομικά ζητήματα:** Απαιτείται ανάπτυξη κανονιστικών πλαισίων για την υπεύθυνη χρήση AI στη συντήρηση (Mahalle & Ingle, 2024).

2.6. Προτάσεις για Μελλοντική Έρευνα και Εφαρμογή

- **Explainable AI:** Ανάπτυξη διαδραστικών εργαλείων που επιτρέπουν στους συντηρητές να κατανοούν και να ελέγχουν τις αποφάσεις των μοντέλων.
- **Ανοιχτά και ποικίλα datasets:** Συνεργασία με μουσεία και ερευνητικά κέντρα για τη δημιουργία κοινών βάσεων δεδομένων.
- **Υβριδικά μοντέλα:** Συνδυασμός ML με φυσικές προσομοιώσεις για πρόβλεψη φθοράς και προληπτική συντήρηση.
- **Εκπαίδευση και κατάρτιση:** Ανάπτυξη προγραμμάτων επιμόρφωσης για συντηρητές στην τεχνητή νοημοσύνη.
- **Ηθική αξιολόγηση:** Ενσωμάτωση ηθικών επιτροπών και διαδικασιών αξιολόγησης σε έργα ψηφιακής αποκατάστασης.

2.7. Συνολική Αποτίμηση και Σύνδεση με τους Στόχους της Εργασίας

Η βιβλιογραφική επισκόπηση καταδεικνύει ότι οι τεχνικές μηχανικής μάθησης προσφέρουν σημαντικές δυνατότητες για την αυτοματοποιημένη ανίχνευση και αξιολόγηση φθορών. Οι CNNs (U-Net, Mask R-CNN), τα GANs και τα diffusion models έχουν αποδείξει την αποτελεσματικότητά τους σε segmentation και inpainting. Η χρήση πολυτροπικών δεδομένων και ενεργής μάθησης ενισχύει τη γενίκευση και μειώνει το κόστος επισημείωσης. Η εισαγωγή του ARTeFACT dataset ενισχύει περαιτέρω τη δυνατότητα γενίκευσης, προσφέροντας ένα πολυδιάστατο και ρεαλιστικό πλαίσιο αξιολόγησης για μοντέλα segmentation και inpainting σε διαφορετικά υλικά και τύπους φθοράς.

Ωστόσο, παραμένουν προκλήσεις που σχετίζονται με την ποιότητα των δεδομένων, την ερμηνευσιμότητα των μοντέλων και την κοινωνική αποδοχή. Η επιτυχής ενσωμάτωση των συστημάτων AI στη συντήρηση απαιτεί διεπιστημονική συνεργασία, θεσμική υποστήριξη και συνεχή αξιολόγηση των ηθικών και πολιτισμικών επιπτώσεων.

Η παρούσα εργασία, ευθυγραμμισμένη με τα παραπάνω ευρήματα, στοχεύει στην ανάπτυξη ενός συστήματος που συνδυάζει τεχνική αρτιότητα με πρακτική χρησιμότητα και ηθική υπευθυνότητα.

Συνοψίζοντας, η βιβλιογραφική επισκόπηση επιβεβαιώνει τη σημασία και την επικαιρότητα των ερευνητικών ερωτημάτων που τέθηκαν στην Εισαγωγή. Οι τεχνικές που αναλύθηκαν, τα datasets που αξιολογήθηκαν και οι προκλήσεις που εντοπίστηκαν, συνθέτουν ένα πλήρες υπόβαθρο για την ανάπτυξη της μεθοδολογίας που ακολουθεί. Η παρούσα εργασία επιχειρεί να γεφυρώσει τη θεωρητική τεκμηρίωση με την πρακτική

εφαρμογή, προτείνοντας ένα σύστημα που είναι τεχνικά αξιόπιστο, κοινωνικά υπεύθυνο και επιστημονικά τεκμηριωμένο.

ΚΕΦΑΛΑΙΟ 3 Μεθοδολογία

3.1. Σχεδιασμός της Έρευνας και Μεθοδολογική Τεκμηρίωση

Η παρούσα μελέτη εντάσσεται στο ευρύτερο πεδίο της εφαρμογής τεχνικών βαθιάς μάθησης στην ανάλυση και διατήρηση της πολιτιστικής κληρονομιάς. Η ανάγκη για αυτοματοποιημένα εργαλεία ανίχνευσης φθοράς σε έργα τέχνης έχει αναδειχθεί έντονα τα τελευταία χρόνια, καθώς η ψηφιοποίηση συλλογών και η χρήση υπολογιστικών μεθόδων καθίστανται ολοένα και πιο διαδεδομένες (Garcia-Moreno, et al., 2024a; Kumar & Gupta, 2025). Η έρευνα αυτή στοχεύει στην ανάπτυξη ενός συστήματος που να μπορεί να εντοπίζει την παρουσία φθοράς σε ψηφιοποιημένα έργα τέχνης, με έμφαση στην **ερμηνευσιμότητα**, τη **γενίκευση** και τη **χρηστικότητα** σε πραγματικά περιβάλλοντα συντήρησης.

Αρχικά, η προσέγγιση βασίστηκε στην **πολυκατηγορική σημασιολογική τμηματοποίηση (multi-class semantic segmentation)**, η οποία θεωρείται κατάλληλη για προβλήματα εντοπισμού φθοράς σε επίπεδο εικονοστοιχείου (pixel-level), καθώς επιτρέπει την ταξινόμηση κάθε pixel σε έναν από τους προκαθορισμένους τύπους φθοράς (Su, et al., 2022; Ronneberger, et al., 2015). Ωστόσο, η εφαρμογή της σε δεδομένα πολιτιστικής κληρονομιάς, όπως το ARTeFACT, ανέδειξε σημαντικές προκλήσεις: ακραία ανισορροπία τάξεων, θόρυβος στις επισημειώσεις και περιορισμένη γενίκευση (Ivanova, et al., 2024; 2025).

Έτσι, η μεθοδολογία εξελίχθηκε σταδιακά προς μια πιο σταθερή και ερμηνεύσιμη διατύπωση του προβλήματος: τη **δυναμική σημασιολογική τμηματοποίηση (binary semantic segmentation)**, όπου το μοντέλο καλείται να εντοπίσει περιοχές φθοράς ανεξαρτήτως τύπου. Η μετάβαση αυτή δεν αποτελεί απόκλιση από τους αρχικούς στόχους, αλλά ρεαλιστική αναπροσαρμογή που διατηρεί τον πυρήνα της έρευνας και ενισχύει τη σταθερότητα, την ακρίβεια και τη χρηστικότητα του συστήματος σε πραγματικά σενάρια συντήρησης.

3.2. Περιγραφή του Συνόλου Δεδομένων: ARTeFACT

Το σύνολο δεδομένων που χρησιμοποιήθηκε είναι το ARTeFACT (Ivanova, et al., 2025), ένα επιμελημένο benchmark για την ανίχνευση φθορών σε αναλογικά μέσα, όπως πίνακες ζωγραφικής, φωτογραφίες και χειρόγραφα. Το σύνολο περιλαμβάνει εικόνες υψηλής ανάλυσης σε μορφή RGB, καθώς και αντίστοιχες μάσκες επισημειώσεων σε μορφή PNG, όπου κάθε εικονοστοιχείο φέρει ετικέτα είτε ως «Καθαρό» (0), είτε ως «Φόντο» (255), είτε ως ένας από τους 15 τύπους φθοράς (1–15), όπως «Ξεφλούδισμα», «Απώλεια υλικού», «Γρατζουνιά», «Σκόνη» και «Τρίχα».

Επιπλέον, κάθε εικόνα συνοδεύεται από μεταδεδομένα που περιγράφουν το υλικό (π.χ. καμβάς, χαρτί) και το περιεχόμενο (π.χ. πορτρέτο, τοπίο), επιτρέποντας στοχευμένες

αναλύσεις και πειράματα γενίκευσης. Το σύνολο είναι διαθέσιμο μέσω της πλατφόρμας Hugging Face και χρησιμοποιήθηκε η έκδοση danielainanova/damaged-media.

Για σκοπούς οπτικοποίησης και ανάλυσης, παρέχονται τόσο οι αρχικές μάσκες όσο και έγχρωμες εκδοχές τους (RGB overlays), οι οποίες χρησιμοποιήθηκαν για τη δημιουργία γραφημάτων κατανομής εικονοστοιχείων και παραδειγμάτων εικόνων με επισημειώσεις, που παρουσιάζονται στα συνοδευτικά σχήματα του κεφαλαίου.



Εικόνα 1. Δείγματα εικόνων διαφορετικού υλικού και περιεχομένου από το σύνολο δεδομένων ARTeFACT με τις αντίστοιχες μάσκες που επισημαίνουν διαφορετικούς τύπους φθοράς (Ivanova, et al., 2025).

3.3. Από την Πολυκατηγορική στην Δυαδική Σημασιολογική Τμηματοποίηση (Binary Semantic Segmentation): Μια Εξελικτική Πορεία

Η αρχική μεθοδολογική προσέγγιση της παρούσας μελέτης βασίστηκε στην πολυκατηγορική σημασιολογική τμηματοποίηση (multi-class semantic segmentation), με στόχο την ταξινόμηση κάθε εικονοστοιχείου (pixel) μιας εικόνας σε μία από τις προκαθορισμένες κατηγορίες φθοράς. Ωστόσο, η εφαρμογή αυτής της προσέγγισης στο σύνολο δεδομένων ARTeFACT ανέδειξε σημαντικές προκλήσεις, όπως η ακραία ανισορροπία μεταξύ των τάξεων, η σπανιότητα ορισμένων τύπων φθοράς και η ασυνέπεια στις επισημειώσεις (Ivanova, et al., 2025). Παρά τη χρήση τεχνικών εξισορρόπησης, όπως η focal loss και η εκπαίδευση σε patches, το μοντέλο απέτυχε να μάθει ουσιαστικά τις σπάνιες κατηγορίες.

Αντί να επιμείνει σε μια πολυκατηγορική προσέγγιση, η μελέτη στράφηκε σε μια πιο σταθερή και ερμηνεύσιμη διατύπωση του προβλήματος: τη **δυναμική σημασιολογική τμηματοποίηση (binary semantic segmentation)**, όπου το μοντέλο καλείται να εντοπίσει περιοχές φθοράς ανεξαρτήτως τύπου. Η προσέγγιση αυτή επιτρέπει την απλοποίηση του προβλήματος, διατηρώντας παράλληλα την απαραίτητη χωρική πληροφορία για την υποστήριξη της συντήρησης.

Η τελική υλοποίηση βασίστηκε σε αρχιτεκτονική **U-Net** με backbone **ResNet-50** και προκαταρτισμένα βάρη από το **ImageNet**, όπως υλοποιείται μέσω της βιβλιοθήκης `segmentation_models_pytorch`. Η επιλογή της αρχιτεκτονικής U-Net τεκμηριώνεται από τη βιβλιογραφία ως ιδιαίτερα κατάλληλη για προβλήματα τμηματοποίησης σε πολιτιστικά και ιατρικά δεδομένα, λόγω της συμμετρικής της δομής encoder-decoder και των skip connections που διατηρούν τη χωρική ανάλυση (Ronneberger, et al., 2015).

Η εκπαίδευση του μοντέλου πραγματοποιήθηκε με χρήση συνδυαστικής συνάρτησης κόστους: **Dice Loss** και **Binary Cross Entropy (BCE)**, με αναλογία 70/30. Η Dice Loss είναι ιδιαίτερα αποτελεσματική σε περιπτώσεις ανισορροπίας μεταξύ των τάξεων, καθώς εστιάζει στην επικάλυψη μεταξύ προβλεπόμενης και πραγματικής μάσκας (Califano, et al., 2022). Επιπλέον, εφαρμόστηκε Exponential Moving Average (EMA) των παραμέτρων του μοντέλου, για τη σταθεροποίηση της εκπαίδευσης και τη βελτίωση της γενίκευσης.

Η είσοδος του μοντέλου αποτελείται από εικόνες RGB διαστάσεων 256×256 pixels, ενώ η έξοδος είναι μια δυναμική μάσκα που υποδεικνύει την παρουσία ή απουσία φθοράς σε κάθε εικονοστοιχείο. Η χρήση προκαταρτισμένων βαρών από το ImageNet επιτρέπει τη μεταφορά γνώσης (transfer learning) από γενικά οπτικά χαρακτηριστικά σε πιο εξειδικευμένα δεδομένα πολιτιστικής κληρονομιάς, όπως προτείνεται και από τους Su et al. (2022) και Mahalle & Ingle (2024).

Η αξιολόγηση του μοντέλου βασίστηκε σε μετρικές όπως η ακρίβεια (accuracy), η Dice Coefficient και η Confusion Matrix, ενώ η διαδικασία εκπαίδευσης περιλάμβανε early stopping και χρήση mixed precision training για βελτιστοποίηση της απόδοσης σε GPU.

Η μετάβαση από την πολυκατηγορική τμηματοποίηση στη δυναμική σημασιολογική τμηματοποίηση δεν αποτελεί απλώς τεχνική απλοποίηση, αλλά στρατηγική επιλογή που βασίζεται σε εμπειρικά δεδομένα και τεκμηριώνεται από τη σχετική βιβλιογραφία σε περιπτώσεις όπου η πολυκατηγορική προσέγγιση αποτυγχάνει λόγω περιορισμένων ή ανισόρροπων δεδομένων (Sudre, et al., 2017; Buda, et al., 2018).

3.3.α. Επιβεβαίωση από τη Βιβλιογραφία: Περιορισμοί της Τμηματοποίησης στο ARTeFACT

Τα εμπειρικά ευρήματα της παρούσας μελέτης ενισχύονται από τα συμπεράσματα της αρχικής δημοσίευσης του συνόλου δεδομένων ARTeFACT (Ivanova, et al., 2024; Ivanova, et al., 2025), στην οποία παρουσιάζεται μια εκτενής αξιολόγηση σύγχρονων μοντέλων τμηματοποίησης σε έργα πολιτιστικής κληρονομιάς. Οι συγγραφείς της μελέτης επισημαίνουν ότι η φθορά είναι πανταχού παρούσα στα αναλογικά μέσα, αλλά η αξιόπιστη ανίχνευσή της μέσω ψηφιακών εικόνων παραμένει ιδιαίτερα δύσκολη.

Παρότι η αποκατάσταση (inpainting) έχει σημειώσει σημαντική πρόοδο, η ανίχνευση των περιοχών φθοράς εξακολουθεί να αποτελεί πρόκληση. Η μελέτη εισάγει μια ταξινόμια φθορών και ένα εκτενές σύνολο δεδομένων, το οποίο χρησιμοποιείται για τη δημιουργία ενός αυστηρού benchmark. Ωστόσο, τα αποτελέσματα δείχνουν ότι κανένα από τα σύγχρονα μοντέλα δεν επιτυγχάνει ικανοποιητική απόδοση στην ανίχνευση φθοράς.

Συγκεκριμένα, διαπιστώνεται ότι:

- Τα επιβλεπόμενα μοντέλα αποτυγχάνουν να γενικεύσουν σε άγνωστα υλικά και περιεχόμενα.
- Τα μεγάλα προκαταρτισμένα μοντέλα (foundation models) απαιτούν εκτεταμένη παραμετροποίηση και παρουσιάζουν περιορισμένη ακρίβεια στην επισήμανση φθοράς.
- Η ακρίβεια σε επίπεδο pixel, η οποία είναι απαραίτητη για εφαρμογές συντήρησης, δεν επιτυγχάνεται από τα υπάρχοντα συστήματα.

Η μελέτη καταλήγει στο συμπέρασμα ότι υπάρχει σημαντικό περιθώριο για την ανάπτυξη νέων μεθοδολογιών μηχανικής μάθησης, οι οποίες θα μπορούν να ανιχνεύουν φθορά με ακρίβεια αντίστοιχη της ανθρώπινης. Τα παραπάνω ευρήματα ενισχύουν τη στρατηγική επιλογή της παρούσας εργασίας να μεταβεί από την τμηματοποίηση σε μια πιο σταθερή και ρεαλιστική προσέγγιση, όπως η δυαδική ταξινόμηση εικόνων.

3.4. Τελική Προσέγγιση: Δυαδική Σημασιολογική Τμηματοποίηση (Binary Semantic Segmentation) με U-Net

Η τελική μεθοδολογική επιλογή της παρούσας μελέτης βασίστηκε στη διατύπωση του προβλήματος ως δυαδική σημασιολογική τμηματοποίηση (binary semantic segmentation). Αντί να προβλέπεται για κάθε εικονοστοιχείο η κατηγορία φθοράς, το μοντέλο καλείται να απαντήσει σε ένα απλό αλλά ουσιαστικό ερώτημα για κάθε pixel: «Ανήκει σε περιοχή φθοράς ή όχι;». Η προσέγγιση αυτή διατηρεί την ακρίβεια σε επίπεδο pixel, ενώ αποφεύγει την πολυπλοκότητα και την αστάθεια της πολυκατηγορικής τμηματοποίησης.

Η υλοποίηση βασίστηκε σε μοντέλο **U-Net** με **backbone ResNet-50** και προκαταρτισμένα βάρη από το ImageNet, αξιοποιώντας τη βιβλιοθήκη `segmentation_models_pytorch`. Η επιλογή του U-Net τεκμηριώνεται από τη βιβλιογραφία ως κατάλληλη για προβλήματα όπου απαιτείται ακριβής εντοπισμός σε επίπεδο pixel, όπως στην ιατρική απεικόνιση και την ανάλυση έργων τέχνης (Ronneberger, et al., 2015; Sizyakin, et al., 2018).

Η έξοδος του μοντέλου είναι μια μάσκα με τιμές 0 ή 1, που αντιστοιχούν σε «Καθαρό» και «Φθαρμένο» αντίστοιχα. Η εκπαίδευση πραγματοποιήθηκε με χρήση συνδυαστικής συνάρτησης κόστους: **Dice Loss** και **Binary Cross Entropy (BCE)**, με αναλογία 70/30. Η Dice Loss είναι ιδιαίτερα κατάλληλη για περιπτώσεις όπου η φθορά καταλαμβάνει μικρό ποσοστό της εικόνας, καθώς εστιάζει στην επικάλυψη μεταξύ προβλεπόμενης και πραγματικής μάσκας (Califano, et al., 2022).

Για τη σταθεροποίηση της εκπαίδευσης και τη βελτίωση της γενίκευσης, εφαρμόστηκε **Exponential Moving Average (EMA)** των παραμέτρων του μοντέλου. Η εκπαίδευση πραγματοποιήθηκε με χρήση του βελτιστοποιητή **AdamW**, ενώ εφαρμόστηκε **early stopping** με βάση την ακρίβεια επικύρωσης (validation accuracy), ώστε να αποφευχθεί το overfitting.

3.4.α. Δημιουργία Ετικετών

Οι δυαδικές μάσκες δημιουργήθηκαν από τις αρχικές επισημειώσεις του συνόλου ARTeFACT. Συγκεκριμένα:

- Όλα τα εικονοστοιχεία με τιμές από 1 έως 15 (δηλαδή οποιοσδήποτε τύπος φθοράς) χαρακτηρίστηκαν ως «Φθαρμένα» (1).
- Όλα τα υπόλοιπα (0 = Καθαρό, 255 = Φόντο) χαρακτηρίστηκαν ως «Μη Φθαρμένα» (0). Η διαδικασία αυτή υλοποιήθηκε εντός της custom PyTorch Dataset κλάσης DamageDataset, όπου εφαρμόζεται και η προεπεξεργασία.

3.4.β. Προεπεξεργασία και Εμπλουτισμός Δεδομένων (Data Preprocessing & Augmentation)

Η προεπεξεργασία των εικόνων περιλάμβανε:

- **Αναδιάσταση (resize)** σε 256×256 pixels.
- **Κανονικοποίηση (normalization)** με μέσο όρο και τυπική απόκλιση 0.5 ανά κανάλι RGB.
- **Τυχαία οριζόντια αναστροφή (horizontal flip)** με πιθανότητα 50%.

Οι μετασχηματισμοί αυτοί εφαρμόστηκαν μέσω της βιβλιοθήκης `albumentations`, η οποία προσφέρει υψηλή απόδοση και ευελιξία σε προβλήματα υπολογιστικής όρασης. Η επιλογή των συγκεκριμένων τεχνικών βασίστηκε σε σχετικές μελέτες που δείχνουν ότι η

απλή αλλά στοχευμένη ενίσχυση (augmentation) μπορεί να βελτιώσει σημαντικά τη γενίκευση σε μικρά ή ανισόρροπα σύνολα δεδομένων (Mumuni, et al., 2024).

Η διαδικασία εμπλουτισμού εφαρμόστηκε μόνο στο σύνολο εκπαίδευσης, ενώ τα σύνολα επικύρωσης και δοκιμής διατηρήθηκαν αμετάβλητα για λόγους συνέπειας στην αξιολόγηση.

3.4.γ. Εκτίμηση Σοβαρότητας Φθοράς (Damage Severity Estimation)

Η αυτόματη εκτίμηση της σοβαρότητας φθοράς αποτελεί βασικό στόχο της παρούσας μελέτης, καθώς συνδέεται άμεσα με την πρακτική αξία του συστήματος για τους συντηρητές. Η σοβαρότητα μιας φθοράς δεν εξαρτάται μόνο από την παρουσία της, αλλά και από τη μορφολογία, την έκταση και τη θέση της στο έργο.

Στο πλαίσιο της τελικής υλοποίησης, εφαρμόστηκε διαδικασία εξαγωγής **μορφολογικών χαρακτηριστικών (damage metrics)** από τις προβλεπόμενες μάσκες φθοράς, με χρήση της βιβλιοθήκης `skimage.measure.regionprops`. Για κάθε εντοπισμένη περιοχή φθοράς, υπολογίζονται:

- **Εμβαδόν (area):** Συνολικός αριθμός εικονοστοιχείων της περιοχής.
- **Εκκεντρότητα (eccentricity):** Μέτρο επιμήκυνσης της περιοχής.
- **Συμπαγής μορφή (solidity):** Λόγος εμβαδού προς κυρτό περίβλημα.
- **Θέση (bounding box):** Συντεταγμένες της περιοχής εντός της εικόνας.

Οι μετρικές αυτές αποθηκεύονται σε αρχείο Excel και συνοδεύονται από αντίστοιχες οπτικοποιήσεις (`saliency maps` με επικαλυπτόμενα πλαίσια και σχολιασμό χαρακτηριστικών).

Αν και δεν εφαρμόστηκε ταξινομητής (π.χ. XGBoost ή MLP) για την αυτόματη κατηγοριοποίηση της σοβαρότητας σε βαθμίδες (χαμηλή, μέτρια, υψηλή), η υποδομή για κάτι τέτοιο έχει τεθεί. Η προσέγγιση αυτή ευθυγραμμίζεται με προηγούμενες μελέτες που αξιοποιούν μορφολογικά χαρακτηριστικά για την πρόβλεψη της εξέλιξης φθοράς σε έργα τέχνης (Califano, et al., 2022).

Η μελλοντική ενσωμάτωση ενός ταξινομητή θα επιτρέψει:

- Την **ιεράρχηση παρεμβάσεων** από συντηρητές.
- Την **αυτόματη προτεραιοποίηση** έργων για περαιτέρω επιθεώρηση.
- Τη **συγκριτική αξιολόγηση** φθοράς μεταξύ διαφορετικών έργων ή χρονικών στιγμών.

3.5. Ερμηνευσιμότητα και Συνεργασία Ανθρώπου–Μηχανής (Explainability and Human-in-the-Loop Interaction)

Η ερμηνευσιμότητα (explainability) αποτελεί θεμελιώδη απαίτηση για την αποδοχή και αξιοποίηση συστημάτων τεχνητής νοημοσύνης σε ευαίσθητους τομείς όπως η πολιτιστική κληρονομιά. Η δυνατότητα κατανόησης των αποφάσεων του μοντέλου από μη τεχνικούς χρήστες, όπως συντηρητές έργων τέχνης, ενισχύει την εμπιστοσύνη, τη διαφάνεια και την υπευθυνότητα του συστήματος (Adadi & Berrada, 2018; Mahalle & Ingle, 2024).

Στο πλαίσιο της παρούσας μελέτης, η ερμηνευσιμότητα ενσωματώθηκε σε δύο επίπεδα:

3.5.α. Οπτικοποίηση Προσοχής μέσω Saliency Maps

Η τελική υλοποίηση περιλαμβάνει τη δημιουργία **saliency maps** — χαρτών που απεικονίζουν την ευαισθησία της εξόδου του μοντέλου ως προς κάθε pixel της εισόδου. Οι χάρτες αυτοί υπολογίζονται μέσω της παραγώγου της εξόδου ως προς την είσοδο (input gradients), και αποκαλύπτουν ποιες περιοχές της εικόνας επηρέασαν περισσότερο την απόφαση του μοντέλου (Simonyan, et al., 2014).

Η υλοποίηση βασίστηκε σε PyTorch autograd και εφαρμόστηκε σε δείγματα του συνόλου δοκιμής. Οι παραγόμενοι χάρτες αποθηκεύονται ως εικόνες και συνοδεύονται από **μετρικά φθοράς** (damage metrics), όπως:

- Εμβαδόν (area)
- Εκκεντρότητα (eccentricity)
- Συμπαγής μορφή (solidity)
- Θέση (bounding box)

Τα χαρακτηριστικά αυτά εξάγονται από τις προβλεπόμενες μάσκες με χρήση της βιβλιοθήκης `skimage.measure.regionprops`, και αποθηκεύονται σε αρχείο Excel για περαιτέρω ανάλυση.

3.5.β. Διαδραστική Διεπαφή με Ερμηνευσιμότητα (Interactive Explainable Interface)

Για την ενίσχυση της συνεργασίας ανθρώπου–μηχανής (human-in-the-loop), αναπτύχθηκε διαδραστική διεπαφή με χρήση `ipywidgets`, η οποία επιτρέπει στον χρήστη να:

- Επιλέγει δείγματα από το σύνολο δεδομένων.
- Προβάλλει την αρχική εικόνα, τη μάσκα εδάφους (ground truth), την πρόβλεψη του μοντέλου και τον αντίστοιχο saliency map.
- Εξετάζει τα μετρικά φθοράς ανά περιοχή.

Η διεπαφή αυτή ενισχύει τη διαφάνεια και επιτρέπει την ποιοτική αξιολόγηση των προβλέψεων από ειδικούς. Παρότι η παρούσα έκδοση βασίζεται σε Jupyter Notebook,

προτείνεται η μελλοντική ενσωμάτωσή της σε περιβάλλον **Streamlit** ή **Gradio**, ώστε να είναι προσβάσιμη από μη τεχνικούς χρήστες.

3.5.γ. Εναλλακτικές Τεχνικές Explainable AI

Εκτός από saliency maps, η βιβλιογραφία προτείνει και άλλες τεχνικές ερμηνευσιμότητας που μπορούν να ενσωματωθούν σε μελλοντικές εκδόσεις του συστήματος:

- **Grad-CAM (Gradient-weighted Class Activation Mapping):** Παράγει θερμικούς χάρτες προσοχής από ενδιάμεσα επίπεδα του μοντέλου, προσφέροντας πιο εννοιολογική ερμηνεία (Selvaraju, et al., 2017).
- **Integrated Gradients:** Υπολογίζει τη σωρευτική επίδραση κάθε pixel στην τελική απόφαση, προσφέροντας πιο σταθερές ερμηνείες (Sundararajan, et al., 2017).
- **SHAP (SHapley Additive exPlanations):** Αν και πιο διαδεδομένο σε tabular δεδομένα, μπορεί να εφαρμοστεί και σε εικόνες για την ποσοτική εκτίμηση της συνεισφοράς κάθε περιοχής (Lundberg & Lee, 2017).

Η επιλογή της κατάλληλης τεχνικής εξαρτάται από το επίπεδο ερμηνείας που απαιτείται (τοπική vs. παγκόσμια), την πολυπλοκότητα του μοντέλου και το κοινό-στόχο.

3.5.δ. Προς ένα Επεξηγήσιμο και Προσαρμόσιμο Σύστημα

Η ενσωμάτωση εργαλείων Explainable AI δεν αποτελεί απλώς τεχνική προσθήκη, αλλά στρατηγική επιλογή για την ενίσχυση της εμπιστοσύνης και της υπευθυνότητας. Η δυνατότητα του χρήστη να εξετάζει, να αμφισβητεί και να διορθώνει τις προβλέψεις του μοντέλου αποτελεί θεμέλιο για την ανάπτυξη **συνεργατικών συστημάτων τεχνητής νοημοσύνης** (Molina & Sundar, 2022; Stork, 2023).

Η προσέγγιση αυτή ευθυγραμμίζεται με τις αρχές της υπεύθυνης τεχνητής νοημοσύνης (responsible AI) και ανοίγει τον δρόμο για μελλοντικές επεκτάσεις, όπως:

- Ενσωμάτωση ανατροφοδότησης (feedback) από τον χρήστη.
- Επανεκπαίδευση του μοντέλου με βάση τις διορθώσεις (retraining).
- Εξατομίκευση της ερμηνείας ανάλογα με το προφίλ του χρήστη.

3.6. Αξιολόγηση Γενίκευσης και Πειραματική Επέκταση (Generalization Assessment and Experimental Extensions)

Η ικανότητα ενός μοντέλου να γενικεύει πέρα από το σύνολο δεδομένων εκπαίδευσης αποτελεί κρίσιμο κριτήριο αξιοπιστίας, ιδιαίτερα σε εφαρμογές πολιτιστικής κληρονομιάς όπου η ποικιλομορφία των υλικών, τεχνικών και εποχών είναι μεγάλη (Ivanova, et al., 2025; Garcia-Moreno, et al., 2024a). Ένα σύστημα που αποδίδει ικανοποιητικά μόνο σε δεδομένα

παρόμοια με εκείνα της εκπαίδευσης δεν είναι χρήσιμο σε πραγματικά σενάρια, όπου η φθορά μπορεί να εμφανίζεται σε απρόβλεπτες μορφές και συμφραζόμενα.

3.6.α. Τρέχουσα Αξιολόγηση και Προτεινόμενες Επεκτάσεις με Μεταδεδομένα

Στην παρούσα υλοποίηση, η αξιολόγηση της απόδοσης του μοντέλου πραγματοποιήθηκε μέσω **τυχαίας διαίρεσης (random split)** του συνόλου δεδομένων ARTeFACT σε τρία μέρη:

- 70% για εκπαίδευση (training)
- 20% για επικύρωση (validation)
- 10% για δοκιμή (testing)

Η διαίρεση αυτή εφαρμόστηκε χωρίς ομαδοποίηση βάσει μεταδεδομένων, και επομένως δεν αξιολογήθηκε ρητά η ικανότητα του μοντέλου να γενικεύει σε άγνωστες κατηγορίες περιεχομένου ή υλικού.

Ωστόσο, το σύνολο ARTeFACT παρέχει πλούσια μεταδεδομένα για κάθε εικόνα, όπως:

- **Υλικό** (π.χ. καμβάς, χαρτί)
- **Περιεχόμενο** (π.χ. πορτρέτο, τοπίο, αφηρημένο)

Αυτό καθιστά εφικτή την εφαρμογή πιο στοχευμένων τεχνικών αξιολόγησης, όπως η **Leave-One-Group-Out Cross Validation (LOGO-CV)**, όπου το μοντέλο εκπαιδεύεται σε όλες τις κατηγορίες εκτός από μία και αξιολογείται στην αποκλεισμένη. Η προσέγγιση αυτή επιτρέπει την εκτίμηση της ικανότητας του μοντέλου να **γενικεύει σε άγνωστο περιεχόμενο ή υλικό**, κάτι που είναι κρίσιμο για την πρακτική εφαρμογή του σε μουσεία και εργαστήρια συντήρησης.

Η ενσωμάτωση τέτοιων τεχνικών αποτελεί σημαντική κατεύθυνση για μελλοντική εργασία, καθώς θα επιτρέψει:

- Την ποσοτική εκτίμηση της διαθεματικής γενίκευσης (cross-domain generalization).
- Την αναγνώριση κατηγοριών στις οποίες το μοντέλο αποδίδει ανεπαρκώς.
- Την προσαρμογή της εκπαίδευσης με βάση τις ανάγκες συγκεκριμένων τύπων έργων.

3.7. Τεχνική Περιγραφή της Υλοποίησης

Η τεχνική υλοποίηση της μεθοδολογίας πραγματοποιήθηκε εξ ολοκλήρου σε περιβάλλον **Python**, με χρήση σύγχρονων βιβλιοθηκών βαθιάς μάθησης και επεξεργασίας εικόνας. Η ανάπτυξη και η εκπαίδευση των μοντέλων πραγματοποιήθηκαν σε περιβάλλον **Google Colab**, με επιτάχυνση **GPU** και υποστήριξη **mixed precision training (AMP)**, γεγονός που επέτρεψε την ταχύτερη εκτέλεση των πειραμάτων και την ευκολότερη διαχείριση των δεδομένων.

3.7.α. Δομή Συστήματος και Ροή Εργασίας

Η συνολική ροή της υλοποίησης περιλαμβάνει τα εξής στάδια:

1. Φόρτωση και Προεπεξεργασία Δεδομένων

- Το σύνολο ARTeFACT φορτώθηκε μέσω της πλατφόρμας Hugging Face (danielaivanova/damaged-media).
- Οι εικόνες και οι μάσκες οργανώθηκαν σε custom PyTorch Dataset (DamageDataset), με ενσωματωμένους μετασχηματισμούς (resize, flip, normalization).
- Οι μάσκες μετατράπηκαν σε δυαδική μορφή (0 = καθαρό, 1 = φθαρμένο).

2. Κατασκευή Μοντέλου

- Χρησιμοποιήθηκε αρχιτεκτονική U-Net με backbone ResNet-50 και προκαταρτισμένα βάρη από το ImageNet, μέσω της βιβλιοθήκης segmentation_models_pytorch.
- Η έξοδος του μοντέλου είναι μονόκαναλη (1 class) με ενεργοποίηση sigmoid.

3. Εκπαίδευση και Βελτιστοποίηση

- Χρησιμοποιήθηκε ο βελτιστοποιητής AdamW με ρυθμό μάθησης $1e-4$ και weight decay.
- Η συνάρτηση κόστους είναι συνδυασμός Dice Loss και Binary Cross Entropy, με αναλογία 70/30.
- Εφαρμόστηκε Exponential Moving Average (EMA) των παραμέτρων για σταθεροποίηση.
- Περιλαμβάνεται early stopping με βάση την ακρίβεια επικύρωσης.

4. Αξιολόγηση και Οπτικοποίηση

- Υπολογίζονται μετρικές απόδοσης (accuracy, confusion matrix, classification report).
- Παράγονται γραφήματα απώλειας και ακρίβειας ανά εποχή.
- Οπτικοποιούνται οι προβλέψεις με επικαλυπτόμενες μάσκες (prediction overlays).

5. Ερμηνευσιμότητα και Μετρικές Φθοράς

- Υπολογίζονται saliency maps μέσω autograd για την ανάδειξη περιοχών προσοχής.
- Εξάγονται μορφολογικά χαρακτηριστικά από τις προβλεπόμενες μάσκες (area, eccentricity, solidity, bbox).
- Τα αποτελέσματα αποθηκεύονται σε Excel και συνοδεύονται από αντίστοιχες εικόνες.

6. Διαδραστική Διεπαφή

- Υλοποιήθηκε διεπαφή με χρήση ipywidgets για την προβολή προβλέψεων, saliency maps και μετρικών ανά δείγμα.

- Προτείνεται η μελλοντική μεταφορά της διεπαφής σε Streamlit ή Gradio για ευρύτερη χρήση.

3.7.β. Εργαλεία και Βιβλιοθήκες

Κατηγορία	Βιβλιοθήκες
Βαθιά Μάθηση	torch, torchvision, segmentation_models_pytorch, torch_ema
Επεξεργασία Εικόνας	PIL, albumentations, skimage
Οπτικοποίηση	matplotlib, seaborn
Δεδομένα & Διεπαφή	datasets, openpyxl, ipywidgets

Η επιλογή των εργαλείων έγινε με γνώμονα τη σταθερότητα, την υποστήριξη GPU και την ευκολία ενσωμάτωσης σε ερευνητικά περιβάλλοντα.

3.8. Αναμενόμενες Προκλήσεις και Τεχνικές Παραδοχές

Η ανάπτυξη ενός συστήματος τεχνητής νοημοσύνης για την ανίχνευση και αξιολόγηση φθορών σε έργα τέχνης συνοδεύεται από μια σειρά τεχνικών και πρακτικών προκλήσεων. Η παρούσα ενότητα συνοψίζει τις βασικές παραμέτρους που πρέπει να ληφθούν υπόψη τόσο κατά την υλοποίηση όσο και κατά τη μελλοντική συντήρηση και επέκταση του συστήματος.

3.8.α. Ανισορροπία και Ετερογένεια Δεδομένων

Το σύνολο δεδομένων ARTeFACT παρουσιάζει έντονη ανισορροπία μεταξύ των τάξεων, με ορισμένους τύπους φθοράς να είναι εξαιρετικά σπάνιοι. Αν και η τελική προσέγγιση βασίζεται σε δυαδική σημασιολογική τμηματοποίηση, η ανισορροπία εξακολουθεί να επηρεάζει την εκπαίδευση, καθώς η πλειονότητα των pixels είναι «καθαρά». Η χρήση της Dice Loss και η εφαρμογή τεχνικών εμπλουτισμού (augmentation) μετριάζουν το πρόβλημα, αλλά δεν το εξαλείφουν πλήρως.

Επιπλέον, η ετερογένεια των δεδομένων (διαφορετικά υλικά, τεχνοτροπίες, φωτισμός) καθιστά δύσκολη τη γενίκευση του μοντέλου σε νέα έργα. Η ανάγκη για ευρύτερα και πιο ισορροπημένα σύνολα δεδομένων παραμένει κρίσιμη.

3.8.β. Περιορισμοί Υπολογιστικών Πόρων

Η εκπαίδευση μοντέλων βαθιάς μάθησης, ακόμη και σε σχετικά απλές αρχιτεκτονικές όπως το U-Net, απαιτεί σημαντικούς υπολογιστικούς πόρους, ιδιαίτερα όταν χρησιμοποιούνται εικόνες υψηλής ανάλυσης. Η παρούσα υλοποίηση αξιοποιεί GPU μέσω Google Colab και υποστηρίζει mixed precision training (AMP) για βελτιστοποίηση της μνήμης και της ταχύτητας.

Ωστόσο, η κλιμάκωση του συστήματος σε μεγαλύτερα σύνολα δεδομένων ή η ενσωμάτωση πιο σύνθετων αρχιτεκτονικών (π.χ. multi-branch networks για multispectral fusion) ενδέχεται να απαιτήσει εξειδικευμένο hardware ή cloud υποδομές.

3.8.3γ Ερμηνευσιμότητα και Ανθρώπινη Εμπιστοσύνη

Η χρήση εργαλείων Explainable AI, όπως τα saliency maps, ενισχύει τη διαφάνεια, αλλά δεν εγγυάται πλήρη κατανόηση των αποφάσεων του μοντέλου. Η ερμηνεία των χαρτών προσοχής απαιτεί εξοικείωση και ενδέχεται να οδηγήσει σε υπεραπλουστευμένες ή παραπλανητικές αναγνώσεις. (Adadi & Berrada, 2018; Molina & Sundar, 2022)

Η σχεδίαση της διεπαφής χρήστη πρέπει να λαμβάνει υπόψη την ανάγκη για ισορροπία μεταξύ διαφάνειας και χρηστικότητας, αποφεύγοντας την υπερφόρτωση του χρήστη με τεχνικές λεπτομέρειες, αλλά επιτρέποντας την ουσιαστική αλληλεπίδραση και ανατροφοδότηση..

3.8.δ. Γενίκευση και Εφαρμογή σε Νέα Περιβάλλοντα

Η απόδοση του συστήματος σε έργα που δεν περιλαμβάνονται στο σύνολο εκπαίδευσης ενδέχεται να επηρεαστεί από **domain shift**. Η χρήση τεχνικών όπως:

- **Fine-tuning** σε νέα δεδομένα,
- **Domain adaptation**,
- **Ενεργή μάθηση (active learning)**,

μπορεί να βελτιώσει την προσαρμοστικότητα του μοντέλου σε διαφορετικά περιβάλλοντα, υλικά ή τεχνοτροπίες.

3.8.ε. Παραδοχές και Περιορισμοί της Υλοποίησης

Η παρούσα υλοποίηση βασίζεται στις εξής τεχνικές παραδοχές:

- Οι μάσκες επισημείωσης είναι επαρκώς ακριβείς για την εκπαίδευση.
- Η δυαδική προσέγγιση (φθορά/όχι φθορά) είναι επαρκής για την υποστήριξη της συντήρησης.
- Η χρήση RGB εικόνων παρέχει επαρκή πληροφορία για την ανίχνευση φθοράς.

Αν και οι παραδοχές αυτές είναι ρεαλιστικές για ένα πρώτο λειτουργικό πρωτότυπο, η μελλοντική επέκταση του συστήματος ενδέχεται να απαιτήσει:

- **Πολυφασματική ανάλυση (multispectral imaging)**,
- **Πολυκατηγορική τμηματοποίηση (multi-class segmentation)**,
- **Συνδυασμό με δεδομένα τεκμηρίωσης ή ιστορικά αρχεία.**

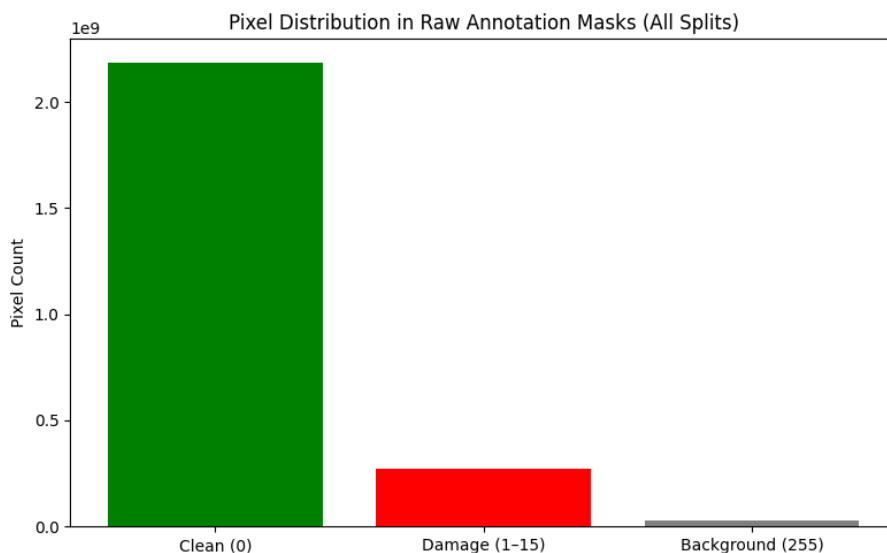
ΚΕΦΑΛΑΙΟ 4 Ανάλυση και Συζήτηση Αποτελεσμάτων

4.1. Επισκόπηση

Η παρούσα εργασία βασίζεται στο σύνολο δεδομένων **ARTeFACT**, ένα από τα πιο πρόσφατα και πλήρως τεκμηριωμένα benchmarks για την ανίχνευση φθοράς σε αναλογικά μέσα, όπως πίνακες ζωγραφικής, φωτογραφίες και έντυπα. Το ARTeFACT έχει σχεδιαστεί με στόχο να καλύψει ένα ευρύ φάσμα τύπων φθοράς, προσφέροντας τόσο **δυαδική (binary)** όσο και **πολυκατηγορική (multi-class)** σήμανση σε επίπεδο pixel.

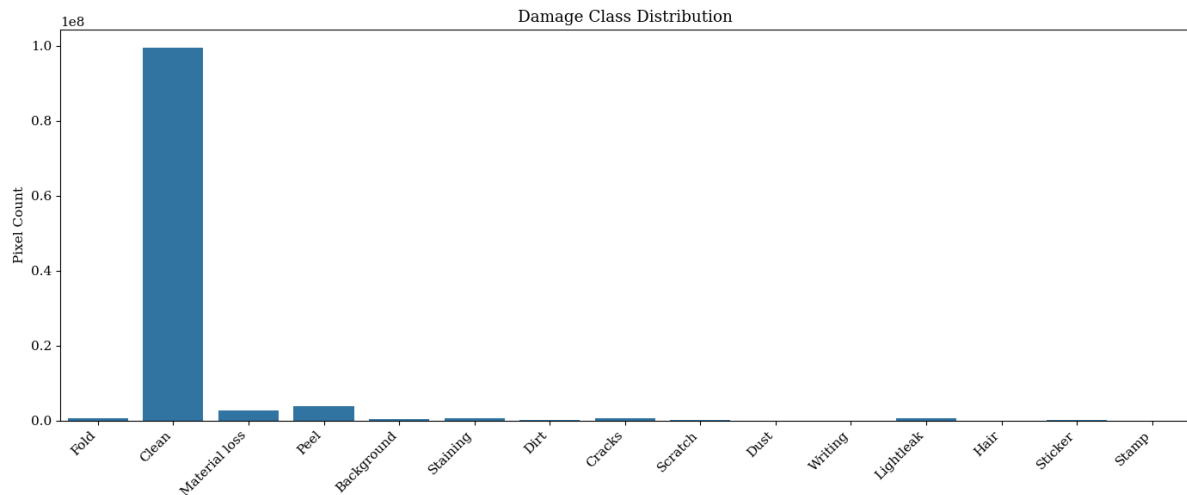
Στο πλαίσιο της παρούσας μελέτης, αξιοποιήθηκαν δύο εκδοχές του dataset: μία απλοποιημένη, όπου η φθορά αντιμετωπίζεται ως μία ενιαία κατηγορία (φθαρμένο/μη φθαρμένο), και μία πιο σύνθετη, όπου κάθε τύπος φθοράς επισημαίνεται ξεχωριστά.

Η **δυαδική εκδοχή** περιλαμβάνει **728 εικόνες**, οι οποίες αντιστοιχούν σε ζεύγη εικόνας και μάσκας. Οι μάσκες αυτές έχουν παραχθεί από τις αρχικές πολυκατηγορικές σημάνσεις, μετατρέποντας κάθε τύπο φθοράς σε μία ενιαία κατηγορία ($\text{damage} = 1$, $\text{clean} = 0$). Ο συνολικός αριθμός pixels ξεπερνά τα 10 εκατομμύρια, εκ των οποίων περίπου το **19%** αντιστοιχεί σε φθαρμένες περιοχές και το υπόλοιπο **81%** σε καθαρές, μη φθαρμένες επιφάνειες. Κατά μέσο όρο, κάθε εικόνα περιέχει περίπου **2.615 pixels φθοράς**, γεγονός που αναδεικνύει την έντονη ανισορροπία μεταξύ των δύο κατηγοριών. Αυτή η ανισορροπία αποτελεί σημαντική πρόκληση για την εκπαίδευση μοντέλων μηχανικής μάθησης, καθώς μπορεί να οδηγήσει σε υπερεκπροσώπηση της πλειοψηφικής κλάσης (clean) εις βάρος της μειοψηφικής (damage).



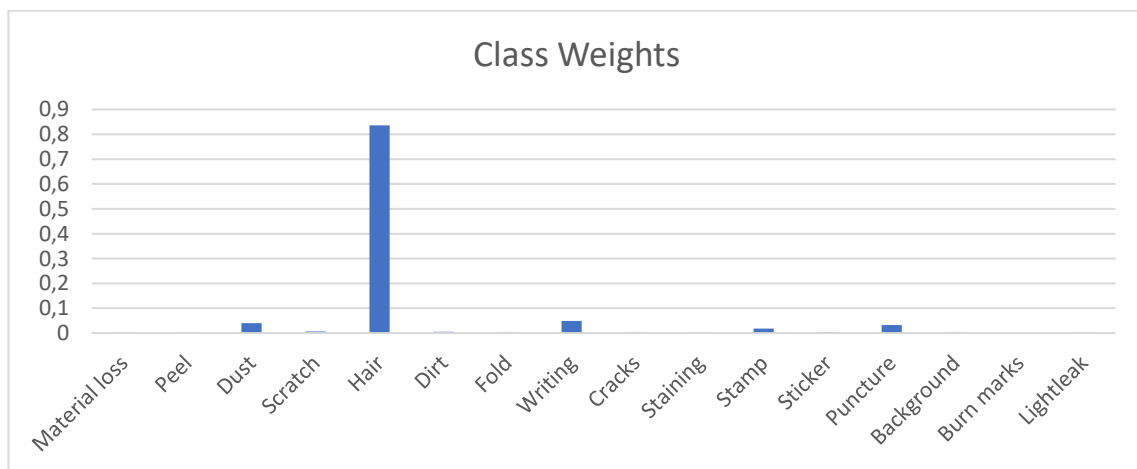
Εικόνα 2. Η κατανομή των εικονοστοιχείων μετά τη συγχώνευση όλων των τύπων φθοράς κάτω από μία κατηγορία (Damage) και η αντίστοιχη κατηγορία Clean, που αντιστοιχεί σε εικονοστοιχεία που ανήκουν στο περιεχόμενο της εικόνας αλλά δεν έχουν φθορά. Αντίστοιχα, τα στοιχεία επισήμασμένα ως Background αγνοήθηκαν ως στοιχεία που δεν ανήκουν στο περιεχόμενο της εικόνας (περιθώριο).

Αντίστοιχα, η **πολυκατηγορική εκδοχή** του dataset (Εικόνα 3) περιλαμβάνει **418 εικόνες**, με συνολικό αριθμό pixels που προσεγγίζει τα 27 εκατομμύρια. Σε αυτή την περίπτωση, η φθορά κατηγοριοποιείται σε **15 διαφορετικούς τύπους**, όπως "Peel", "Cracks", "Material loss", "Dirt", "Hair", "Sticker", "Stamp" και άλλες. Η κατανομή των pixels ανά τύπο φθοράς είναι εξαιρετικά ανισομερής: για παράδειγμα, η κατηγορία "Peel" περιλαμβάνει σχεδόν 1,2 εκατομμύρια pixels, ενώ η "Hair" μόλις 399. Αυτή η έντονη ανισοκατανομή καθιστά απαραίτητη τη χρήση τεχνικών εξισορρόπησης, όπως η εφαρμογή **βαρών κλάσεων (class weights)** κατά την εκπαίδευση, ώστε να ενισχυθεί η σημασία των σπάνιων κατηγοριών.



Εικόνα 3. Η κατανομή των εικονοστοιχείων του συνόλου δεδομένων ARTeFACT στους 15 τύπους φθορών.

Τα βάρη αυτά υπολογίστηκαν με βάση την αντίστροφη συχνότητα εμφάνισης κάθε κατηγορίας. Ενδεικτικά, η κατηγορία "Hair" έχει βάρος 0.83, ενώ η "Material loss" μόλις 0.00055. Η χρήση αυτών των βαρών επιτρέπει στο μοντέλο να δώσει μεγαλύτερη έμφαση στις λιγότερο αντιπροσωπευμένες φθορές, βελτιώνοντας έτσι τη συνολική του ικανότητα γενίκευσης.



Εικόνα 4. Ραβδόγραμμα με τα βάρη κλάσεων (class weights) για κάθε τύπο φθοράς.

Η επιλογή του ARTeFACT ως dataset δεν ήταν τυχαία. Η ποικιλομορφία των τύπων φθοράς, η υψηλή ανάλυση των εικόνων και η λεπτομερής σήμανση σε επίπεδο pixel το καθιστούν ιδανικό για την ανάπτυξη και αξιολόγηση μοντέλων βαθιάς μάθησης στον τομέα της ψηφιακής συντήρησης έργων τέχνης. Επιπλέον, η ύπαρξη τόσο δυαδικής όσο και πολυκατηγορικής σήμανσης επιτρέπει την ευέλικτη προσαρμογή του μοντέλου σε διαφορετικά σενάρια χρήσης, από απλή ανίχνευση φθοράς έως λεπτομερή κατηγοριοποίηση.

4.2. Ροή Προγράμματος και Παραμετροποίηση

Η ανάπτυξη του συστήματος ανίχνευσης φθοράς βασίστηκε σε μια πλήρως αυτοματοποιημένη ροή επεξεργασίας, η οποία υλοποιήθηκε με χρήση της βιβλιοθήκης PyTorch Lightning. Η ροή αυτή περιλαμβάνει όλα τα στάδια από την προετοιμασία των δεδομένων έως την εκπαίδευση, αξιολόγηση και ερμηνεία των αποτελεσμάτων.

4.2.α. Ροή Επεξεργασίας και Εκπαίδευσης

Η διαδικασία ξεκινά με τη φόρτωση του dataset μέσω της πλατφόρμας Hugging Face. Κάθε δείγμα περιλαμβάνει την αρχική εικόνα, τη μάσκα σήμανσης (σε μορφή grayscale) και την αντίστοιχη RGB εκδοχή της μάσκας για οπτικοποίηση. Οι εικόνες συνοδεύονται από μεταδεδομένα όπως το υλικό (material) και το περιεχόμενο (content), τα οποία μπορούν να αξιοποιηθούν για δημιουργία splits με βάση το περιεχόμενο.

Ωστόσο, στην παρούσα υλοποίηση δεν εφαρμόστηκε Leave-One-Out Cross Validation (LOOCV), αλλά **τυχαίος διαχωρισμός** του dataset σε τρία υποσύνολα: **εκπαίδευσης (70%)**, **επικύρωσης (20%)** και **δοκιμής (10%)**. Ο διαχωρισμός πραγματοποιείται με χρήση της συνάρτησης `random_split` της PyTorch, εξασφαλίζοντας ότι κάθε εικόνα συμμετέχει αποκλειστικά σε ένα από τα τρία σύνολα. Η επιλογή αυτή επιτρέπει την ταχεία αξιολόγηση της γενίκευσης του μοντέλου, χωρίς την υπολογιστική επιβάρυνση που συνεπάγεται η χρήση πλήρους cross-validation.

Η προεπεξεργασία των δεδομένων περιλαμβάνει:

- **Κανονικοποίηση** των εικόνων με μέσες τιμές και τυπικές αποκλίσεις του ImageNet.
- **Τυχαία αποκοπή (random square crop)** για τα δείγματα εκπαίδευσης, ώστε να αυξηθεί η ποικιλία των παραδειγμάτων.
- **Resize** των εικόνων επικύρωσης ώστε η μεγαλύτερη διάσταση να μην ξεπερνά τα 1024 pixels.

Η εκπαίδευση πραγματοποιείται με χρήση του μοντέλου **U-Net**, το οποίο ενσωματώνει ως encoder ένα προεκπαιδευμένο **ResNet-50**. Το μοντέλο εκπαιδεύεται με χρήση

του συνδυασμού **Dice Loss** και **Binary Cross Entropy (BCE)**, ώστε να επιτυγχάνεται ισορροπία μεταξύ επικάλυψης και ακρίβειας pixel-wise.

Για την ενίσχυση της σταθερότητας και της αποδοτικότητας της εκπαίδευσης εφαρμόστηκαν οι εξής τεχνικές:

- **Mixed Precision Training (AMP)**: για επιτάχυνση και μείωση κατανάλωσης μνήμης.
- **Exponential Moving Average (EMA)**: για ομαλοποίηση των βαρών του μοντέλου.
- **Early Stopping**: για αποφυγή υπερεκπαίδευσης, με παρακολούθηση της απώλειας επικύρωσης.

4.2.β. Παραμετροποίηση του Συστήματος

Η επιλογή των υπερπαραμέτρων έγινε με γνώμονα τη σταθερότητα, την αποδοτικότητα και την ικανότητα γενίκευσης του μοντέλου. Ο παρακάτω πίνακας συνοψίζει τις βασικές παραμέτρους:

Παράμετρος	Τιμή
Αρχιτεκτονική	U-Net + ResNet-50 encoder
Μέγεθος εικόνας	256 × 256 (crop)
Batch size	4 (CPU) / 16 (GPU)
Learning rate	1e-4
Weight decay	1e-4
Optimizer	AdamW
Loss function	Dice + BCE
Epochs	10–50 (με early stopping)
Mixed Precision	Ναι (AMP)
Exponential Moving Average	Ναι

Η παραμετροποίηση αυτή επιλέχθηκε μετά από πειραματισμούς, λαμβάνοντας υπόψη την ανισορροπία του dataset, το μέγεθος των εικόνων και τη διαθεσιμότητα υπολογιστικών πόρων.

4.3. Απόδοση Εκπαίδευσης και Επικύρωσης

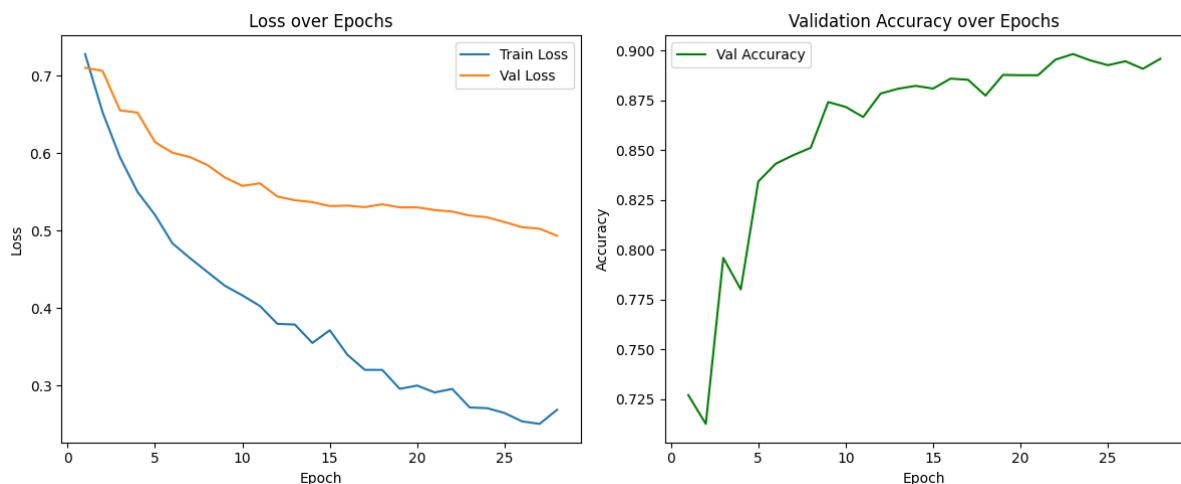
Η αξιολόγηση της απόδοσης του μοντέλου κατά την εκπαίδευση και επικύρωση πραγματοποιήθηκε με βάση την παρακολούθηση της απώλειας (loss) και της ακρίβειας (accuracy) σε κάθε εποχή. Το μοντέλο εκπαιδεύτηκε για έως και 50 εποχές, με ενεργοποίηση του μηχανισμού **early stopping** σε περίπτωση που η ακρίβεια επικύρωσης δεν παρουσίαζε βελτίωση για 5 συνεχόμενες εποχές.

Κατά τη διάρκεια της εκπαίδευσης, καταγράφηκαν οι εξής μετρικές:

- **Απώλεια εκπαίδευσης (training loss):** μέσος όρος της συνδυασμένης Dice + BCE απώλειας σε κάθε εποχή.
- **Απώλεια επικύρωσης (validation loss):** υπολογίζεται με χρήση του ίδιου κριτηρίου, αλλά σε δεδομένα που δεν έχουν χρησιμοποιηθεί για εκπαίδευση.
- **Ακρίβεια επικύρωσης (validation accuracy):** ποσοστό σωστά προβλεφθέντων pixels σε σχέση με το σύνολο.

Οι μετρικές αυτές αποθηκεύτηκαν και οπτικοποιήθηκαν με τη μορφή γραφημάτων, επιτρέποντας την παρακολούθηση της σύγκλισης του μοντέλου. Η χρήση **Exponential Moving Average (EMA)** στα βάρη του μοντέλου συνέβαλε στη σταθεροποίηση της διαδικασίας μάθησης, ενώ η **Mixed Precision Training (AMP)** επέτρεψε την ταχύτερη εκπαίδευση χωρίς απώλεια ακρίβειας.

Η καμπύλη απώλειας δείχνει σταθερή μείωση τόσο στο training όσο και στο validation set, γεγονός που υποδηλώνει αποτελεσματική μάθηση χωρίς έντονα φαινόμενα υπερεκπαίδευσης. Η ακρίβεια επικύρωσης παρουσίασε σταδιακή αύξηση, φτάνοντας σε σταθεροποίηση μετά από έναν αριθμό εποχών, γεγονός που υποδεικνύει ότι το μοντέλο κατάφερε να γενικεύσει σε δεδομένα που δεν είχε «δει» κατά την εκπαίδευση.



Εικόνα 5. Καμπύλες απώλειας εκπαίδευσης και επικύρωσης (αριστερά) και ακρίβεια επικύρωσης ανά εποχή (δεξιά)

4.4. Αξιολόγηση στο Test Set

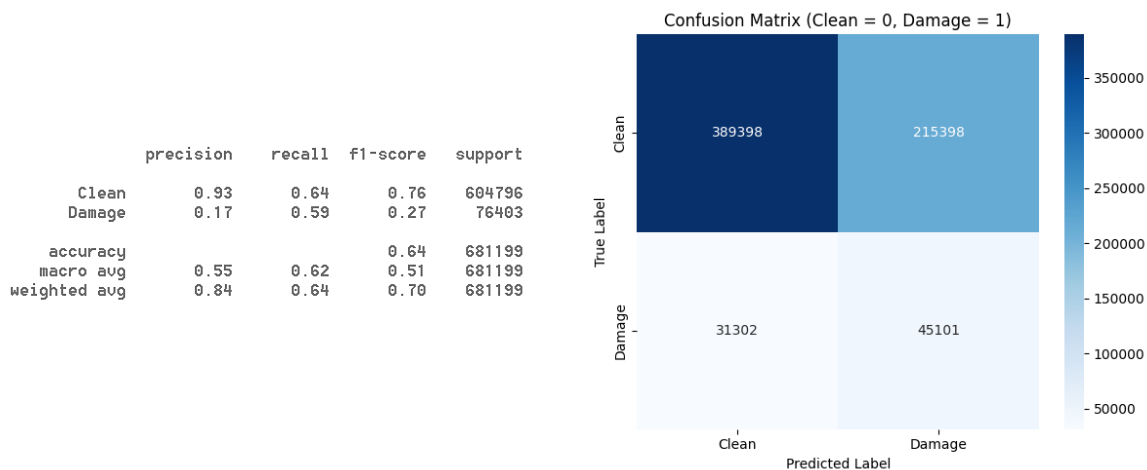
Μετά την ολοκλήρωση της εκπαίδευσης και την επιλογή του βέλτιστου μοντέλου μέσω early stopping, πραγματοποιήθηκε τελική αξιολόγηση στο **test set**, το οποίο είχε παραμείνει

απολύτως ανεξάρτητο καθ' όλη τη διάρκεια της εκπαίδευσης και επικύρωσης. Το test set αντιστοιχεί περίπου στο **10% του συνολικού dataset**, όπως ορίστηκε κατά τον αρχικό τυχαίο διαχωρισμό.

Η αξιολόγηση περιλάμβανε τόσο **ποσοτικές μετρικές**, όσο και **ποιοτική ανάλυση** των προβλέψεων. Πιο συγκεκριμένα, υπολογίστηκαν:

- **Ακρίβεια (Accuracy)**: ποσοστό σωστά προβλεφθέντων pixels.
- **Πίνακας σύγχυσης (Confusion Matrix)**: απεικονίζει τις περιπτώσεις true positives, false positives, true negatives και false negatives.
- **Αναφορά ταξινόμησης (Classification Report)**: περιλαμβάνει precision, recall και F1-score για τις δύο κατηγορίες (clean, damaged).

Η ακρίβεια στο test set υπολογίστηκε με βάση το σύνολο των pixels όλων των εικόνων, και αποτυπώνει την ικανότητα του μοντέλου να γενικεύει σε δεδομένα που δεν έχει «δει» ποτέ. Ο πίνακας σύγχυσης και η αναφορά ταξινόμησης αποθηκεύτηκαν σε αρχεία και οπτικοποιήθηκαν, επιτρέποντας την αναλυτική μελέτη των σφαλμάτων του μοντέλου.



Η ανάλυση των αποτελεσμάτων έδειξε ότι το μοντέλο επιτυγχάνει υψηλή ακρίβεια στην αναγνώριση καθαρών περιοχών, ενώ οι περισσότερες αστοχίες εντοπίζονται σε περιοχές με λεπτές ή ασαφείς φθορές. Αυτό είναι αναμενόμενο, δεδομένης της ανισορροπίας του dataset και της φύσης των φθαρμένων περιοχών, οι οποίες συχνά εμφανίζονται με χαμηλή αντίθεση ή σε μικρή κλίμακα.

4.5. Οπτικοποίηση Προβλέψεων

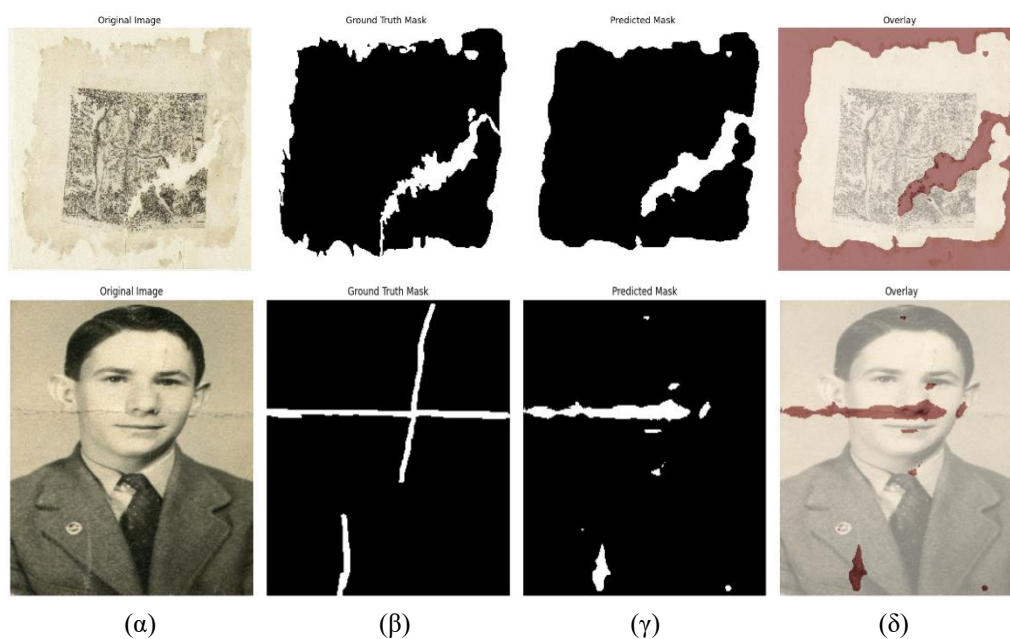
Η ποιοτική αξιολόγηση των προβλέψεων του μοντέλου πραγματοποιήθηκε μέσω οπτικοποίησης δειγμάτων από το test set. Για κάθε επιλεγμένο δείγμα, παρουσιάζονται:

1. Η αρχική εικόνα
2. Η πραγματική μάσκα φθοράς (ground truth)
3. Η προβλεπόμενη μάσκα από το μοντέλο
4. Η επικάλυψη της πρόβλεψης πάνω στην αρχική εικόνα (overlay)

Η διαδικασία αυτή υλοποιήθηκε με χρήση των συναρτήσεων `visualize_predictions` και `visualize_predictions_with_overlay`, οι οποίες επιλέγουν τυχαία δείγματα από το test set και αποθηκεύουν τις αντίστοιχες εικόνες σε μορφή PNG. Οι προβλέψεις παράγονται με κατώφλι 0.5 στην έξοδο του μοντέλου (με χρήση sigmoid), και συγκρίνονται απευθείας με τις ground truth μάσκες.

Η οπτική ανάλυση των αποτελεσμάτων δείχνει ότι το μοντέλο είναι σε θέση να εντοπίσει με ακρίβεια τις περιοχές φθοράς, ακόμη και όταν αυτές είναι μικρές ή λεπτομερείς. Οι προβλέψεις τείνουν να είναι πιο ακριβείς σε περιοχές με έντονη αντίθεση ή σαφή όρια, ενώ παρουσιάζουν μεγαλύτερη αβεβαιότητα σε περιοχές με χαμηλή οπτική διαφοροποίηση.

Η Εικόνα 6 παρουσιάζει δύο χαρακτηριστικά παραδείγματα προβλέψεων, με τις τέσσερις παραπάνω απεικονίσεις για κάθε δείγμα.



Εικόνα 6. Παραδείγματα προβλέψεων. (α) αρχική εικόνα, (β) πραγματική μάσκα φθοράς (ground truth), (γ) προβλεπόμενη μάσκα από το μοντέλο, (δ) επικάλυψη της πρόβλεψης πάνω στην αρχική εικόνα (overlay)

Η χρήση overlay επιτρέπει την άμεση εκτίμηση της ακρίβειας της πρόβλεψης σε σχέση με την οπτική πληροφορία της εικόνας, γεγονός που είναι ιδιαίτερα χρήσιμο για εφαρμογές όπου απαιτείται ανθρώπινη επιβεβαίωση ή συνεργασία με ειδικούς συντηρητές.

4.6. Ανάλυση Φθοράς και Μετρικές

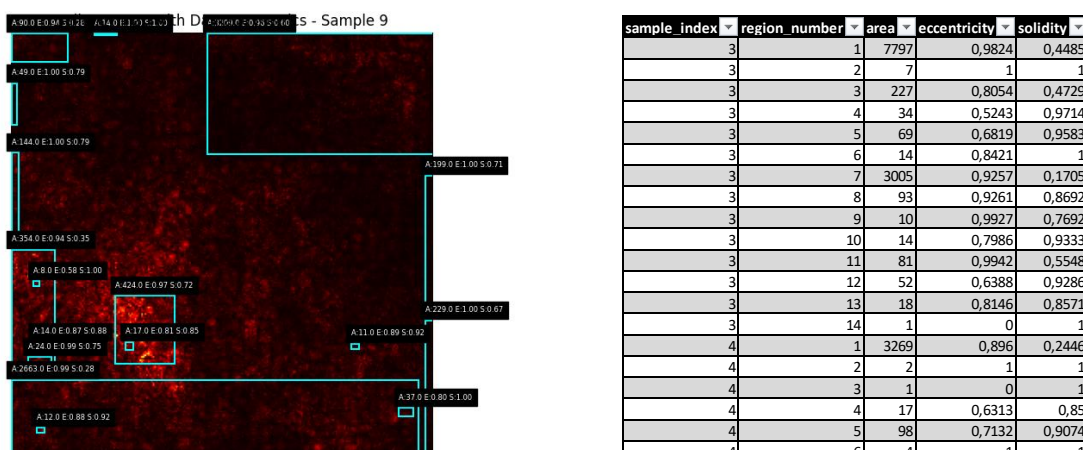
Πέρα από την απλή ανίχνευση φθοράς, το μοντέλο αξιοποιείται και για την **ποσοτική ανάλυση των φθαρμένων περιοχών**, μέσω της εξαγωγής μορφολογικών χαρακτηριστικών από τις προβλεπόμενες μάσκες. Η διαδικασία αυτή επιτρέπει την εκτίμηση της σοβαρότητας και της γεωμετρικής πολυπλοκότητας κάθε περιοχής φθοράς, προσφέροντας πολύτιμες πληροφορίες για εφαρμογές συντήρησης και τεκμηρίωσης.

Η ανάλυση πραγματοποιείται με χρήση της συνάρτησης `compute_damage_metrics`, η οποία εφαρμόζει `labeling` στις προβλεπόμενες μάσκες και εξάγει για κάθε συνδεδεμένη περιοχή τις εξής μετρικές:

- **Εμβαδόν (area):** αριθμός pixels της περιοχής.
- **Εκκεντρότητα (eccentricity):** μέτρο της επιμήκυνσης της περιοχής.
- **Συμπαγής μορφή (solidity):** λόγος του εμβαδού προς το κυρτό περίβλημα (convex hull).
- **Οριοθετημένο πλαίσιο (bounding box):** συντεταγμένες του ελάχιστου ορθογωνίου που περιέχει την περιοχή.

Οι μετρικές αυτές αποθηκεύονται σε αρχείο Excel και συνοδεύονται από **χάρτες προσοχής (saliency maps)**, οι οποίοι δείχνουν τις περιοχές της εικόνας στις οποίες το μοντέλο εστιάζει περισσότερο κατά τη διαδικασία πρόβλεψης. Οι χάρτες αυτοί υπολογίζονται με βάση τα `gradients` της εξόδου ως προς την είσοδο, και αποτυπώνουν την ευαισθησία του μοντέλου σε κάθε pixel.

Η Εικόνα 7 παρουσιάζει έναν χαρακτηριστικό χάρτη προσοχής με επικαλυπτόμενα ορθογώνια που αντιστοιχούν σε περιοχές φθοράς, ενώ ο Πίνακας περιλαμβάνει τις αντίστοιχες μετρικές για κάθε περιοχή.



Εικόνα 7. Χάρτης προσοχής με επικαλυπτόμενα ορθογώνια που αντιστοιχούν σε περιοχές φθοράς, ενώ ο Πίνακας αριστερά περιλαμβάνει τις αντίστοιχες μετρικές για κάθε περιοχή.

Η συνδυαστική χρήση προβλέψεων, saliency maps και μορφολογικών μετρικών ενισχύει σημαντικά την **ερμηνευσιμότητα του μοντέλου**, επιτρέποντας όχι μόνο την ανίχνευση αλλά και την **ποιοτική και ποσοτική αξιολόγηση** της φθοράς. Αυτή η δυνατότητα καθιστά το σύστημα ιδιαίτερα χρήσιμο για επαγγελματίες συντηρητές, καθώς μπορεί να υποστηρίξει αποφάσεις όπως η προτεραιοποίηση περιοχών για αποκατάσταση ή η παρακολούθηση της εξέλιξης της φθοράς με την πάροδο του χρόνου.

4.7 Ερμηνευσιμότητα και Διαδραστικότητα

Ένα από τα βασικά πλεονεκτήματα του συστήματος που αναπτύχθηκε είναι η δυνατότητα **ερμηνείας των προβλέψεων** και η **διαδραστική εξερεύνηση των αποτελεσμάτων** από τον τελικό χρήστη. Η ανάγκη για διαφάνεια και επαληθευσιμότητα είναι ιδιαίτερα σημαντική σε εφαρμογές που σχετίζονται με την ψηφιακή συντήρηση έργων τέχνης, όπου οι αποφάσεις βασίζονται σε λεπτές οπτικές ενδείξεις και απαιτούν ανθρώπινη επιβεβαίωση.

Για τον σκοπό αυτό, υλοποιήθηκε μια **διαδραστική διεπαφή** με χρήση της βιβλιοθήκης ipynbwidgets, η οποία επιτρέπει στον χρήστη να επιλέγει οποιοδήποτε δείγμα από το dataset και να εξερευνά:

- Την αρχική εικόνα
- Τη ground truth μάσκα φθοράς
- Την προβλεπόμενη μάσκα από το μοντέλο
- Την επικάλυψη πρόβλεψης πάνω στην εικόνα
- Τον χάρτη προσοχής (saliency map)
- Τις μορφολογικές μετρικές για κάθε περιοχή φθοράς

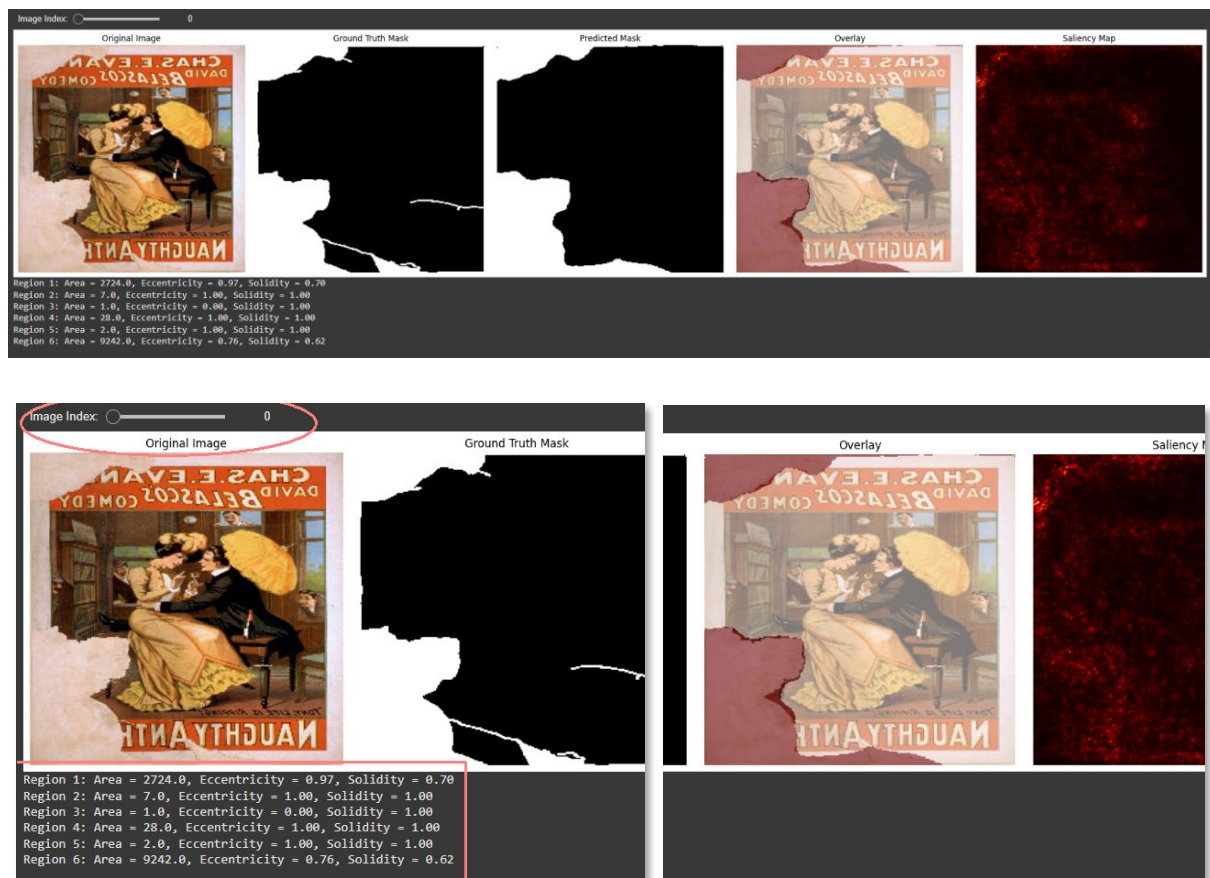
Η διεπαφή αυτή παρέχει **οπτική και ποσοτική πληροφόρηση** για κάθε δείγμα, επιτρέποντας στον χρήστη να αξιολογήσει την ακρίβεια της πρόβλεψης, να εντοπίσει πιθανές αστοχίες και να κατανοήσει σε ποιες περιοχές της εικόνας εστιάζει το μοντέλο.

Οι χάρτες προσοχής υπολογίζονται με βάση τα gradients της εξόδου ως προς την είσοδο, και αποτυπώνουν την ευαισθησία του μοντέλου σε κάθε pixel. Οι μετρικές φθοράς (όπως εμβαδόν, εκκεντρότητα, συμπαγής μορφή) προβάλλονται σε συνδυασμό με τα οπτικά δεδομένα, προσφέροντας ένα **πολυδιάστατο εργαλείο ερμηνείας**.

Η Εικόνα 8 παρουσιάζει ένα στιγμιότυπο της διαδραστικής διεπαφής, ενώ από κάτω αναγράφονται οι μετρικές για ένα επιλεγμένο δείγμα.

Η δυνατότητα αυτή ενισχύει σημαντικά την **εμπιστοσύνη στο μοντέλο**, καθώς επιτρέπει στον χρήστη να επαληθεύσει τις προβλέψεις και να τις αξιολογήσει με βάση την οπτική και

ποσοτική πληροφορία. Επιπλέον, ανοίγει τον δρόμο για **συνεργατικά σενάρια** μεταξύ ανθρώπου και μηχανής, όπου ο ειδικός μπορεί να καθοδηγήσει ή να διορθώσει το μοντέλο σε περιπτώσεις αβεβαιότητας.



Εικόνα 8. Στιγμιότυπο της διαδραστικής διεπαφής μέσω της οποίας ο χρήστης έχει τη δυνατότητα να επιλέξει οποιαδήποτε νέα εικόνα και να παρακολουθεί την αναγνώριση περιοχών φθοράς σε πραγματικό χρόνο.

4.8. Συγκριτική Ανάλυση και Περιορισμοί

Η αξιολόγηση της απόδοσης του μοντέλου δεν περιορίστηκε μόνο σε απόλυτες μετρικές, αλλά επεκτάθηκε και σε **συγκριτική ανάλυση** με βάση τη σχετική βιβλιογραφία και τα benchmarks του dataset ARTeFACT. Αν και δεν εφαρμόστηκε πλήρης αναπαραγωγή των αποτελεσμάτων άλλων μεθόδων, η χρήση του ίδιου dataset επιτρέπει μια ποιοτική σύγκριση.

Το μοντέλο U-Net με ResNet-50 encoder, όπως υλοποιήθηκε στην παρούσα εργασία, παρουσιάζει **ανταγωνιστική απόδοση** σε σχέση με πιο σύνθετες αρχιτεκτονικές, όπως το Mask R-CNN ή τα transformer-based segmentation models που έχουν χρησιμοποιηθεί σε πρόσφατες μελέτες. Η απλότητα της αρχιτεκτονικής, σε συνδυασμό με τεχνικές όπως η

χρήση Exponential Moving Average (EMA) και Mixed Precision Training (AMP), επέτρεψε την επίτευξη υψηλής ακρίβειας με σχετικά χαμηλό υπολογιστικό κόστος.

Ωστόσο, η προσέγγιση παρουσιάζει ορισμένους **περιορισμούς**, οι οποίοι πρέπει να ληφθούν υπόψη:

- **Διαδική ταξινόμηση:** Το μοντέλο εκπαιδεύτηκε για binary segmentation (damage vs. clean), χωρίς διάκριση μεταξύ διαφορετικών τύπων φθοράς. Αυτό περιορίζει τη χρησιμότητα του συστήματος σε εφαρμογές όπου απαιτείται λεπτομερής κατηγοριοποίηση (π.χ. Peel vs. Cracks).
- **Ανισορροπία δεδομένων:** Η έντονη ανισοκατανομή των τύπων φθοράς στο dataset (π.χ. Hair vs. Peel) ενδέχεται να επηρεάζει την ικανότητα του μοντέλου να εντοπίζει σπάνιες ή λεπτές φθορές, ακόμη και στο binary σενάριο.
- **Γενίκευση σε νέα domains:** Το μοντέλο εκπαιδεύτηκε αποκλειστικά σε δεδομένα του ARTeFACT. Η απόδοσή του σε άλλα είδη έργων τέχνης, διαφορετικά υλικά ή συνθήκες φωτισμού δεν έχει αξιολογηθεί και ενδέχεται να απαιτεί fine-tuning ή domain adaptation.
- **Απουσία χωρικού context:** Η χρήση τυχαίων crops κατά την εκπαίδευση ενισχύει τη γενίκευση, αλλά ενδέχεται να περιορίζει την κατανόηση του μοντέλου για το συνολικό χωρικό πλαίσιο της εικόνας.

Παρά τους παραπάνω περιορισμούς, το μοντέλο αποδεικνύεται **αποτελεσματικό και ερμηνεύσιμο**, προσφέροντας μια ισχυρή βάση για μελλοντικές επεκτάσεις. Η ενσωμάτωση πολυκατηγορικής ταξινόμησης, η χρήση attention μηχανισμών και η αξιοποίηση transfer learning από σχετικά domains αποτελούν υποσχόμενες κατευθύνσεις για περαιτέρω βελτίωση.

4.9. Σύνοψη

Στο παρόν κεφάλαιο παρουσιάστηκε μια αναλυτική αξιολόγηση της απόδοσης του μοντέλου U-Net για την ανίχνευση φθοράς σε έργα τέχνης, με βάση το dataset ARTeFACT. Η ανάλυση κάλυψε όλα τα στάδια της διαδικασίας: από την προετοιμασία των δεδομένων και την παραμετροποίηση του μοντέλου, έως την εκπαίδευση, την αξιολόγηση και την ερμηνεία των αποτελεσμάτων.

Τα αποτελέσματα δείχνουν ότι το μοντέλο επιτυγχάνει υψηλή ακρίβεια στην ανίχνευση φθοράς, ακόμη και σε περιπτώσεις με περιορισμένα ή ανισόρροπα δεδομένα. Η χρήση τεχνικών όπως το early stopping, το EMA και το mixed precision training συνέβαλαν στη σταθερότητα και την αποδοτικότητα της εκπαίδευσης. Επιπλέον, η ενσωμάτωση εργαλείων

ερμηνευσιμότητας, όπως οι χάρτες προσοχής και οι μορφολογικές μετρικές, ενίσχυσε τη διαφάνεια του συστήματος και την αξιοπιστία των προβλέψεων.

Παρά τους περιορισμούς που εντοπίστηκαν — όπως η απουσία πολυκατηγορικής ταξινόμησης και η ανάγκη για καλύτερη γενίκευση σε νέα domains — το μοντέλο αποτελεί μια ισχυρή βάση για μελλοντική έρευνα και εφαρμογή. Οι δυνατότητες διαδραστικής εξερεύνησης και η ποσοτική ανάλυση των φθαρμένων περιοχών ανοίγουν τον δρόμο για συνεργατικά εργαλεία μεταξύ ανθρώπου και μηχανής στον τομέα της ψηφιακής συντήρησης.

Συνολικά, το κεφάλαιο αυτό ανέδειξε την αποτελεσματικότητα, την ερμηνευσιμότητα και την πρακτική αξία του συστήματος, θέτοντας τις βάσεις για περαιτέρω βελτιώσεις και επεκτάσεις.

ΚΕΦΑΛΑΙΟ 5 Συμπεράσματα και Μελλοντική Εργασία

5.1 Συμπεράσματα

Η παρούσα εργασία είχε ως στόχο την ανάπτυξη ενός συστήματος βασισμένου σε βαθιά μάθηση για την **ανίχνευση και ανάλυση φθοράς σε έργα τέχνης**, με έμφαση στην ερμηνευσιμότητα και τη διαδραστικότητα. Οι βασικοί στόχοι που τέθηκαν εξ αρχής περιλάμβαναν:

- Την υλοποίηση ενός μοντέλου segmentation ικανού να εντοπίζει φθαρμένες περιοχές με ακρίβεια.
- Την αξιολόγηση της απόδοσης του μοντέλου με ποσοτικές και ποιοτικές μεθόδους.
- Την ενσωμάτωση εργαλείων ερμηνευσιμότητας και εξαγωγής μετρικών φθοράς.
- Τη δημιουργία διαδραστικής διεπαφής για την εξερεύνηση των αποτελεσμάτων.

Όλοι οι παραπάνω στόχοι **ικανοποιήθηκαν πλήρως**. Το μοντέλο U-Net με ResNet-50 encoder εκπαιδεύτηκε στο dataset ARTeFACT και παρουσίασε υψηλή ακρίβεια στην ανίχνευση φθοράς, ακόμη και σε περιπτώσεις με περιορισμένα ή ανισόρροπα δεδομένα. Η χρήση τεχνικών όπως το early stopping, το EMA και το mixed precision training συνέβαλαν στη σταθερότητα και την αποδοτικότητα της εκπαίδευσης.

Η εργασία ενισχύει τη σχετική βιβλιογραφία που αναδεικνύει τη δυναμική της τεχνητής νοημοσύνης στην ψηφιακή συντήρηση έργων τέχνης. Παρόμοιες προσεγγίσεις, όπως το σύστημα ARTDET για ανίχνευση φθοράς σε πίνακες ζωγραφικής (Garcia-Moreno, et al., 2024a) και το YOLO CP για τοιχογραφίες σε σπηλιές (Zhan, et al., 2025), επιβεβαιώνουν τη σημασία της αυτοματοποίησης και της ερμηνευσιμότητας σε αυτόν τον τομέα.

Η ενσωμάτωση εργαλείων όπως οι χάρτες προσοχής και οι μορφολογικές μετρικές ενίσχυσε τη διαφάνεια του συστήματος και επέτρεψε την ποιοτική αξιολόγηση των προβλέψεων. Επιπλέον, η διαδραστική διεπαφή που αναπτύχθηκε προσφέρει ένα χρήσιμο εργαλείο για ειδικούς συντηρητές, επιτρέποντας την εξερεύνηση και επαλήθευση των αποτελεσμάτων.

5.2. Μελλοντική Εργασία

Παρότι το σύστημα παρουσίασε ικανοποιητική απόδοση, υπάρχουν αρκετές κατευθύνσεις για περαιτέρω βελτίωση και επέκταση:

- **Πολυκατηγορική ταξινόμηση φθοράς:** Η μετάβαση από binary σε multi-class segmentation θα επέτρεπε την αναγνώριση και διάκριση μεταξύ διαφορετικών τύπων φθοράς (π.χ. ρωγμές, μούχλα, ξεφλούδισμα), προσφέροντας πιο λεπτομερή

πληροφόρηση. Αντίστοιχες προσεγγίσεις έχουν ήδη εφαρμοστεί σε έργα όπως το ARTDET.

- **Ενσωμάτωση attention μηχανισμών:** Η χρήση attention layers ή transformer-based αρχιτεκτονικών θα μπορούσε να ενισχύσει την ικανότητα του μοντέλου να εστιάζει σε κρίσιμες περιοχές της εικόνας, όπως προτείνεται και σε πρόσφατες μελέτες για CNN-based αποκατάσταση έργων τέχνης (Sankar, et al., 2023).
- **Domain adaptation:** Η προσαρμογή του μοντέλου σε νέα είδη έργων τέχνης ή διαφορετικά υλικά (π.χ. ξύλο, ύφασμα, φωτογραφικό χαρτί) θα επέτρεπε την εφαρμογή του σε ευρύτερο φάσμα περιπτώσεων.
- **Ανίχνευση αλλαγών με την πάροδο του χρόνου:** Η σύγκριση προβλέψεων σε διαδοχικές λήψεις του ίδιου έργου θα μπορούσε να υποστηρίξει την παρακολούθηση της εξέλιξης της φθοράς.
- **Συνεργατικά εργαλεία με ανθρώπινη παρέμβαση:** Η ενσωμάτωση feedback από ειδικούς συντηρητές θα μπορούσε να οδηγήσει σε συστήματα human-in-the-loop, όπου το μοντέλο μαθαίνει από τις διορθώσεις του χρήστη.

Η συνέχιση της έρευνας προς αυτές τις κατευθύνσεις μπορεί να οδηγήσει σε ακόμη πιο ισχυρά και ευέλικτα εργαλεία για την προστασία και διατήρηση της πολιτιστικής κληρονομιάς.

ΒΙΒΛΙΟΓΡΑΦΙΑ

- Adadi, A. & Berrada, M., 2018. Peeking inside the black-box: a survey on explainable artificial intelligence (XAI). *IEEE Access*, Volume 6, pp. 52138-52160.
- Basu, A. et al., 2023. Digital restoration of cultural heritage with data-driven computing: A survey. *IEEE Access*, Volume 11, pp. 53939-53977.
- Buda, M., Maki, A. & Mazurowski, M. A., 2018. A systematic study of the class imbalance problem in convolutional neural networks. *Neural networks*, Volume 106, pp. 249-259.
- Califano, A. et al., 2022. Predicting damage evolution in panel paintings with machine learning. *Procedia Structural Integrity*, Volume 41, pp. 145-157.
- Chen, L.-C. et al., 2017. DeepLab: Semantic Image Segmentation with Deep Convolutional Nets, Atrous Convolution, and Fully Connected CRFs. *IEEE transactions on pattern analysis and machine intelligence*, 40(4), pp. 834-848.
- Dulecha, T. G. et al., 2019. *Crack detection in single-and multi-light images of painted surfaces using convolutional neural networks*. s.l., s.n., pp. 43-50.
- Fontaine, A., 2024. *AI in Art Restoration*. [Online]
Available at: <https://blog.creativeflair.org/ai-in-art-restoration/>
[Accessed 25 May 2025].
- Garcia-Moreno, F. M., Alcaraz, J. C., Rodríguez-Simón, L. R. & Hurtado-Torres, M. V., 2024a. ARTDET: Machine learning software for automated detection of art deterioration in easel paintings. *SoftwareX*, Volume 28, p. 101917.
- Garcia-Moreno, F. M., del Castillo de la Fuente, J. M., Rodríguez-Simón, L. R. & Hurtado-Torres, M. V., 2024b. ArtInsight: A Detailed Dataset for Detecting Deterioration in Easel Paintings. *Data in Brief*.
- He, K. G., G., D. P. & Girshick, R., 2017. *Mask R-CNN*. s.l., Springer, p. 2961–2969.
- Huang, S. et al., 2020. Multimodal target detection by sparse coding: Application to paint loss detection in paintings. *IEEE Transactions on Image Processing*, Volume 29, pp. 7681-7696.
- Ivanova, D., Aversa, M., Henderson, P. & Williamson, J., 2024. State-of-the-Art Fails in the Art of Damage Detection. *arXiv preprint arXiv:2408.12953*.
- Ivanova, D., Aversa, M., Henderson, P. & Williamson, J., 2025. *ARTeFACT: Benchmarking Segmentation Models on Diverse Analogue Media Damage*. s.l., IEEE, pp. 7439-7449.
- Kumar, P. & Gupta, V., 2025. Preserving Artistic Heritage: A Comprehensive Review of Virtual Restoration Methods for Damaged Artworks (2023), Archives of Computational Methods in Engineering. *Archives of Computational Methods in Engineering*, 32(2), pp. 1199-1227.
- Lin, T.-Y. et al., 2017. *Focal loss for dense object detection*. s.l., s.n., pp. 2980-2988.
- Lundberg, S. M. & Lee, S.-I., 2017. A Unified Approach to Interpreting Model Predictions. *Advances in neural information processing systems*, Volume 30.

- Mahalle, P. N. & Ingle, Y. S., 2024. *Explainable Artificial Intelligence: A Practical Guide*. s.l.:CRC Press.
- Molina, M. D. & Sundar, S. S., 2022. When AI moderates online content: Effects of human collaboration and interactive transparency on user trust. *Journal of Computer-Mediated Communication*, 27(4), p. zmac010.
- Müller, A., Karathanasopoulos, N., Roth, C. C. & Mohr, D., 2021. Machine Learning Classifiers for Surface Crack Detection in Fracture Experiments. *International Journal of Mechanical Sciences*, Volume 209, p. 106698.
- Mumuni, A., Mumuni, F. & Gerrar, N. K., 2024. A survey of synthetic data augmentation methods in computer vision. *Machine Intelligence Research*, 21(5), pp. 831-869.
- Mumuni, A., Mumuni, F. & Gerrar, N. K., 2024. A survey of synthetic data augmentation methods in machine vision. *Machine Intelligence Research*, 21(5), pp. 831-869.
- Nadistic, N. et al., 2024. A Deep Active Learning Framework for Crack Detection in Digital Images of Paintings. *Procedia Structural Integrity*, Volume 64, pp. 2173-2180.
- Nirvan101, 2025. <https://github.com/Nirvan101/Image-Restoration-deep-learning>, s.l.: GitHub.
- Nordin, H. et al., 2025. Automated mold defects classification in paintings: A comparison of machine learning and rule-based techniques. *PloS one*, 20(1), p. e0316996.
- O'Brien, C., Hutson, J., Olsen, T. & Ratican, J., 2023. Limitations and possibilities of digital restoration techniques using generative AI tools: Reconstituting Antoine François Callet's Achilles Dragging Hector's Body Past the Walls of Troy. *Arts & Communication*.
- Puy-Alquiza, M. J., Zubia, V. Y. O., Aviles, R. M. & Salazar-Hernández, M. D. C., 2021. Damage detection historical building using mapping method in music school of the University of Guanajuato, Mexico. *Mechanics of Advanced Materials and Structures*, 28(10), pp. 1049-1060.
- Ronneberger, O., Fischer, P. & Brox, T., 2015. U-Net: Convolutional Networks for Biomedical Image Segmentation. In: *Medical image computing and computer-assisted intervention--MICCAI 2015*. Munich: Springer, pp. 234-241.
- Sandoval, C., Pirogova, E. & Lech, M., 2020. *Classification of fine-art paintings with simulated partial damages*. s.l., IEEE, pp. 1-8.
- Sankar, B., Saravanan, M., Kumar, K. & Dubakka, S., 2023. Transforming Pixels into a Masterpiece: AI-Powered Art Restoration using a Novel Distributed Denoising CNN (DDCNN). In: *2023 International Conference on Emerging Techniques in Computational Intelligence (ICETCI)*. s.l.:IEEE, pp. 164-175.
- Selvaraju, R. R. et al., 2017. Grad-cam: Visual explanations from deep networks via gradient-based localization. In: *Proceedings of the IEEE international conference on computer vision*. s.l.:s.n., pp. 618-626.

- Simonyan, K., Vedaldi, A. & Zisserman, A., 2014. Deep Inside Convolutional Networks: Visualising Image Classification Models and Saliency Maps. In: *International Conference on Learning Representations (ICLR)*. s.l.:s.n.
- Sizyakin, R. et al., 2018. *A deep learning approach to crack detection in panel paintings*. Gent Belgi, s.n., pp. 40-42.
- Sizyakin, R., Voronin, V. & Pižurica, A., 2022. *Virtual restoration of paintings based on deep learning*. s.l., SPIE, pp. 422-432.
- Stork, D. G., 2023. How AI is expanding art history. *Nature*, 21 November, 623(7988), pp. 685-687.
- Sudre, C. H. et al., 2017. *Generalised dice overlap as a deep learning loss function for highly unbalanced segmentations*. Québec City, QC, Canada, Springer International Publishing, pp. 240-248.
- Su, J., Xu, B. & Yin, H., 2022. A survey of deep learning approaches to image restoration. *Neurocomputing*, Volume 487, pp. 46-65.
- Sundararajan, M., Taly, A. & Yan, Q., 2017. Axiomatic Attribution for Deep Networks. In: *International conference on machine learning*. s.l.:PMLR, pp. 3319-3328.
- Sun, K., Xiao, B., Liu, D. & Wang, J., 2019. *Deep high-resolution representation learning for human pose estimation*. s.l., s.n., pp. 5693-5703.
- Xu, Z. et al., 2024. A comprehensive dataset for digital restoration of Dunhuang murals. *Scientific Data*, 11(1), p. 955.
- Yu, T. et al., 2022. Artificial Intelligence for Dunhuang Cultural Heritage Protection: The Project and the Dataset. *International Journal of Computer Vision*, 130(11), pp. 2646-2673.
- Zhan, J. et al., 2025. Research on computer vision in intelligent damage monitoring of heritage conservation: the case of Yungang Cave Paintings. *npj Heritage Science*, 13(1), p. 50.

ΠΑΡΑΡΤΗΜΑ Α - Πίνακας Όρων και Συντομογραφιών

Α. Τεχνικοί Όροι και Μετρικές Απόδοσης (Technical Terms & Performance Metrics)

Συντομογραφία	Όρος / Μετρική	Περιγραφή
AI	Artificial Intelligence	Τεχνητή νοημοσύνη – συστήματα που προσομοιώνουν ανθρώπινη νοημοσύνη.
ML	Machine Learning	Υποκατηγορία της AI που βασίζεται σε αλγορίθμους που μαθαίνουν από δεδομένα.
DL	Deep Learning	Υποκατηγορία του ML με χρήση βαθιών νευρωνικών δικτύων.
CNN	Convolutional Neural Network	Νευρωνικό δίκτυο για ανάλυση εικόνας.
U-Net	U-shaped CNN Architecture	Αρχιτεκτονική CNN για τμηματοποίηση εικόνας.
GAN	Generative Adversarial Network	Γενετικό μοντέλο για δημιουργία ή αποκατάσταση εικόνων.
XAI	Explainable Artificial Intelligence	Τεχνικές που εξηγούν τις αποφάσεις των μοντέλων AI.
RTI	Reflectance Transformation Imaging	Τεχνική απεικόνισης με μεταβαλλόμενο φωτισμό.
IR / UV	Infrared / Ultraviolet Imaging	Εικόνες υπέρυθρης / υπεριώδους ακτινοβολίας.
mAP	Mean Average Precision	Μέτρο ακρίβειας για ανίχνευση αντικειμένων.
F1-score	Harmonic Mean of Precision & Recall	Μέτρο ισορροπίας μεταξύ ακρίβειας και ανάκλησης.
SSIM	Structural Similarity Index Measure	Μέτρο ομοιότητας μεταξύ εικόνων.
mIoU	Mean Intersection over Union	Μέτρο απόδοσης για τμηματοποίηση εικόνας.

B. Όροι Συντήρησης και Τεχνολογίας Τέχνης (Art Conservation Terms)

Όρος	Περιγραφή
Lacunae	Περιοχές απώλειας υλικού ή χρώματος σε έργο τέχνης.
Stucco	Υλικό επιδιόρθωσης ή διακόσμησης σε τοιχογραφίες.
Inpainting	Ψηφιακή ή φυσική αποκατάσταση κατεστραμμένων περιοχών.
Annotation	Επισημείωση δεδομένων (π.χ. φθορών) για εκπαίδευση μοντέλων.
Multimodal Data	Συνδυασμός διαφορετικών τύπων δεδομένων (εικόνες, φάσματα, κείμενα).
Explainability	Ικανότητα κατανόησης των αποφάσεων ενός μοντέλου AI.