

DETC2025-167593

NEXT-LAYER DEFECTS FORECASTING IN LASER DIRECTED ENERGY DEPOSITION USING CONVOLUTIONAL LONG SHORT-TERM MEMORY AUTOENCODER

Yuxuan Xie^{1,2}, Lequn Chen¹, Youxiang Chew¹, A Senthil Kumar²

¹ Advanced Remanufacturing and Technology Centre (ARTC),
Agency for Science and Technology and Research (A*STAR), 3 CleanTech Loop, #01-01 CleanTech Two,
637143 Singapore

² College of Design and Engineering, Department of Mechanical Engineering,
National University of Singapore, Block EA, #07-08, 9 Engineering Drive 1, 117575 Singapore

ABSTRACT

Laser Direct Energy Deposition (L-DED) is an additive manufacturing technology used to produce large, high-value metal components with complex geometries. However, defects such as cracks, keyhole pores, and uneven surfaces frequently occur, necessitating real-time monitoring and predictive control for quality assurance. While existing optical, acoustic, and multimodal monitoring strategies effectively identify defects, they primarily focus on defect detection rather than predictive forecasting and defect mitigation. To address this limitation, this study proposes a novel framework for next-layer defect forecasting based on historical time-series co-axial melt pool images. The key contribution is the development of a Convolutional Long Short-Term Memory (ConvLSTM) autoencoder model that captures the spatiotemporal evolution of melt pool dynamics to forecast next-layer melt pool images. The forecasted melt pool images are then fed into a Convolutional Neural Network (CNN) model to classify potential defects. The proposed hybrid ConvLSTM-CNN architecture enables proactive defect prediction by generating and classifying next-layer melt pool images. Preliminary results from single-bead wall experiments demonstrate that the ConvLSTM model achieves over 98% pixel-wise accuracy in forecasting future melt pool evolution, while the classifier independently attains over 90% defect classification accuracy. This work lays the foundation for adaptive process adjustments and preemptive defect mitigation in L-DED.

Keywords: Laser Direct Energy Deposition, Process Monitoring, Convolutional Long Short-Term Memory, Defect Prediction, Deep Learning

1. INTRODUCTION

Laser direct energy deposition (L-DED) is a highly efficient additive manufacturing technology used to produce large, complex, and high-value metal components [1]. Despite its numerous advantages, maintaining consistent quality throughout the L-DED process remains a significant challenge due to its complicated defect formation mechanisms and process variability.

In recent years, extensive research has been focused on minimizing inherent flaws in the L-DED process by optimizing critical process parameters, such as laser power, powder flow, carrier gas flow, scanning speed, and feeding rate [2]. However, despite these optimizations, defects such as cracks and keyhole pores still occur due to inherent issues such as material properties, residual stress accumulation, and dynamic deposition process. Consequently, the effective process monitoring is essential for quality control and defect prevention during manufacturing.

Advanced monitoring systems have been developed to detect the occurrence of defects during the manufacturing process. For instance, optical sensors [3-6], including (Charge-coupled device) CCD, high-speed and thermal cameras, enable the analysis of melt pool morphology and thermal history to classify defects. In addition, acoustic sensors detect defects by analyzing acoustic emissions from laser-material interactions [7-9]. Prior research has demonstrated that co-axial melt pool imaging and acoustic emission analysis can support the development of machine learning (ML) models for defect detection and process classification (e.g., non-defective, crack, and keyhole pore regimes) [5], [7], [9-12]. However, the current monitoring approaches largely remain reactive. They typically

intervene only after the defect mechanisms are already initiated. As a result, defect prevention before the formation of defects remains a significant challenge. Therefore, achieving a defect-free L-DED process requires a shift toward proactive strategies. In particular, proactive strategies focus on defect forecasting, where predictive models anticipate defect formation in advance. However, despite promising developments in monitoring and control, current forecasting approaches remain underdeveloped. The gap between monitoring and forecasting underscores the critical need for predictive capabilities and adaptive control mechanisms to mitigate defect formation and improve the process stability.

A review of recent research shows that current monitoring systems in additive manufacturing are primarily focused on detection rather than prediction. For example, Kononenko et al [13] developed an in-situ acoustic emission-based crack detection system for Laser Powder Bed Fusion (LPBF), achieving an accuracy of up to 99% with an inference interval of 1 ms. Similarly, Estalaki et al [14] employed ML models using thermographic data to detect printing states in LPBF with spatial registration, attaining a maximum F1 score of 0.966. In addition, our previous work [9] [15] introduced both acoustic-based and multimodal monitoring systems for defect classification in L-DED. However, the research above remains a reactive stage for defect detection, and the proactive forecasting capabilities are unexplored.

Recently, there have been a few attempts to tackle the challenge of underdeveloped forecasting approaches. For instance, Kim et al [12] proposed a novel melt pool image registration technique to estimate melt pool positions. However, this study focused on estimating the position and connectivity of melt pools, rather than forecasting future printing states. Furthermore, Abranovic et al. [16] introduced a Convolutional Long Short-Term Memory (ConvLSTM)-based approach for predicting printing anomalies, but their work was limited to single-bead printing. Capturing spatiotemporal relationships in multi-layer and complex geometries remains a significant challenge. Similarly, Ko et al [17] proposed a spatiotemporal correlation prediction framework for LPBF processes. While their framework demonstrates promising capabilities, applying it to L-DED is still unexplored due to the different spatiotemporal dynamics involved.

To address limitations in the lack of defect forecasting approaches in L-DED processes, this paper proposes a novel hybrid ConvLSTM-Convolutional Neural Network (CNN) architecture for the proactive next-layer defect forecasting in L-DED processes. The ConvLSTM autoencoder captures the spatiotemporal relationship across layers to generate the next-layer melt pool images from previous layer melt pool images, and the predicted melt pool sequences are then fed to a CNN model to classify potential defects. To validate the proposed approach, preliminary experiments were conducted using co-axial melt pool images from single-bead wall printing processes. This work lays the foundation for real-time adaptive process adjustments and preemptive defect mitigation in L-DED.

The key novelty of this work is the development of the ConvLSTM-CNN model capturing the spatiotemporal evolution for forecasting the next-layer melt pool image and predicting the potential defects. A fundamental challenge in L-DED defect prediction lies in understanding the spatiotemporal relationships within the process. Recent studies indicate that defect formation in L-DED is a progressive phenomenon, transitioning from an initial defect-free state to mild imperfections (e.g., uneven surfaces, cracks) and, in severe cases, keyhole pores [13-15]. These defects are driven by cumulative factors such as dimensional accuracy, heat accumulation, and material instability. More importantly, imperfections in earlier layers influence subsequent deposition tracks and layers, emphasizing the spatiotemporal dependency of defect evolution [18]. Therefore, accurate process modeling and defect prediction necessitates capturing both spatial and temporal aspects of the process.

The remainder of the paper is as follows. Section 2 introduces the novel ConvLSTM-CNN architecture. Section 3 discusses case study and visualization of forecasting melt pool images. Section 4 presents results and discussions. Section 5 concludes this paper with concluding remarks and future work.

2. MATERIALS AND METHODS

2.1 Overview Workflow

The overview of the workflow is illustrated as FIGURE 1 with highlighted contributions. The process monitoring data is the co-axial melt pool video. In this study, co-axial melt pool video is utilized. Video frames are extracted at a rate of 10 Hz, capturing one melt pool image every 100 milliseconds. After cropping, each image is labeled based on its printing state, including non-defective, keyhole pores, laser defocus, and laser-off. The labeled images are then sequenced and organized according to their respective layer and track positions. Each sequence represents the melt pool images for a specific track or layer. These sequences are used as input for the proposed ConvLSTM autoencoder, which is trained to predict melt pool images for the subsequent track or layer. The predictions are then fed into a fine-tuned CNN model, and the model outputs a classification of the predicted images (non-defective, defective, and laser-off), which accomplishes the defect prediction.

2.2 Experiment Setup

Several L-DED experiments using an industrial KUKA robotic arm are conducted. The experimental setup, including the robotic laser module, melt pool position, and monitoring sensors, is shown in FIGURE 2. A co-axial visible spectrum CCD camera is mounted on the laser module for co-axial melt pool imaging. The CCD camera with an acquisition frequency of 30 Hz, recorded the co-axial melt pool evolutions in videos during the L-DED process. One experiment of printing process is fully recorded in one video. In the data pre-process stage, the video is sampled at a frequency of 10 Hz, capturing one data point (or one image) every 100 milliseconds, and every image is labeled with the time stamp according to the video. The total number of melt pool images in single-bead wall printing is 2500. The melt pool images are sequenced according to the time stamp in order

of the track and layer. In the next step, the sequenced images are utilized to train the ConvLSTM autoencoder. Additionally, it is worth noting that the acoustic emission in L-DED experiments is recorded but is not the main topic of this work.

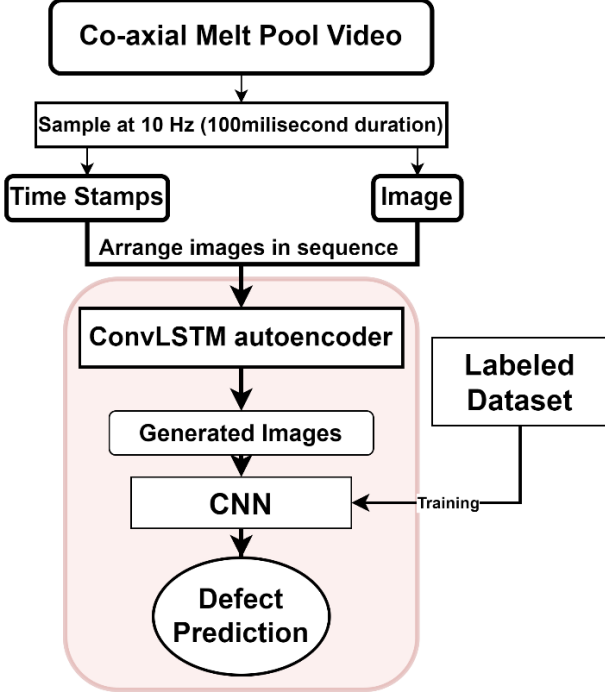


FIGURE 1: OVERVIEW OF THE PROPOSED WORKFLOW OF MONITORING DATA, NEXT-LAYER-PREDICTION BASED ON CONV LSTM AUTOENCODER AND PRE-TRAINED CNN.

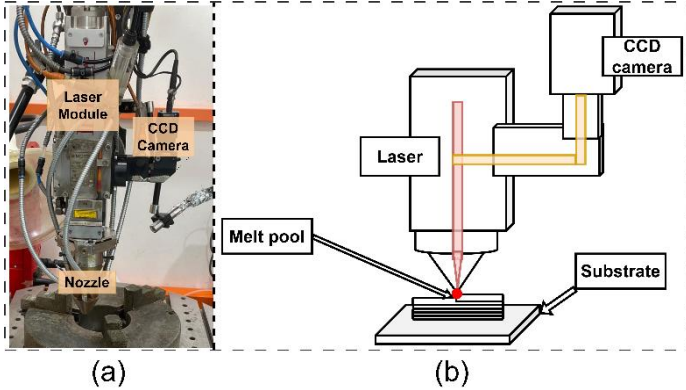


FIGURE 2: (A) INDUSTRIAL LASER-DIRECT ENERGY DEPOSITION SYSTEM INTEGRATED WITH CCD CAMERA, (B) SCHEMATIC VIEW OF THE EXPERIMENT SETUP

2.3 Problem Formulation

The objective of the proposed deep learning model is to predict future co-axial melt pool images based on previously recorded sequences of co-axial melt pool images. Specifically, the model aims to forecast the next sequence of melt pool images with the same length as the input sequence. For the single-bead wall printing process, the input is a sequence of melt pool images from one layer, and the output is the sequence of melt pool

images for the subsequent layer. This relationship is mathematically represented in Equation (1):

$$\sum_{t=0}^i t_k^{j+1} = \sum_{t=0}^i \alpha t_k^j \quad (1)$$

In this equation, t_k^j represents the deposition state at time step t , within layer j and track k . The coefficient α represents the influence of previous layer j on the corresponding points of layer or track $j+1$. In other words, α describes the relationship between tracks (spatiotemporal relationship vector).

For rectangular and square geometries, the printing process becomes more complex due to the influence of adjacent tracks. In track-wise forecasting of multi-track geometries, neighboring tracks from the same and previous layer contribute to the printing outcomes of the next track. The mathematical formulation for this relationship is shown in Equation (2):

$$\sum_{t=0}^i t_k^{j+1} = \sum_{t=0}^i (\alpha t_k^j + f(t_{k-1}^j) + q(t_{k+1}^j)) \quad (2)$$

In this equation, t_k^j represents the printing states from the same track position k in the previous layer j , as the same as in (1), while t_{k-1}^j and t_{k+1}^j represent the printing states of previous and next tracks at previous layer j . The functions $f(\cdot)$ and $q(\cdot)$ represent spatiotemporal relationships and interactions from these two adjacent tracks in the previous layer. This formulation incorporates both temporal and spatial relationships, enabling the ConvLSTM autoencoder to capture complex dependencies within the printing process, particularly for multi-track geometries such as rectangular and square patterns. To illustrate these interactions more clearly, FIGURE 3 demonstrates the spatial relationships between the tracks in the previous and current layers and the interested track being predicted by the ConvLSTM model.

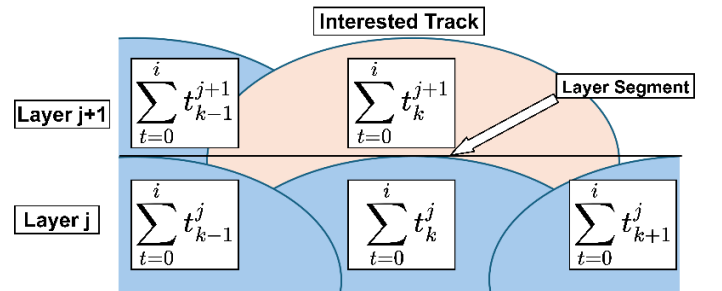


FIGURE 3: SECTION VIEW OF AN INTERESTED TRACK WITH INFLUENCES FROM ADJACENT TRACKS IN THE CURRENT LAYER AND PREVIOUS LAYER.

2.4 Data Preparation

In the experiment designs, six single-bead walls components are printed with process parameters. The dimensions of these printed components and process parameters are listed in TABLE 1. During the printing process, co-axial melt pool videos are captured using the visible spectrum CCD camera.

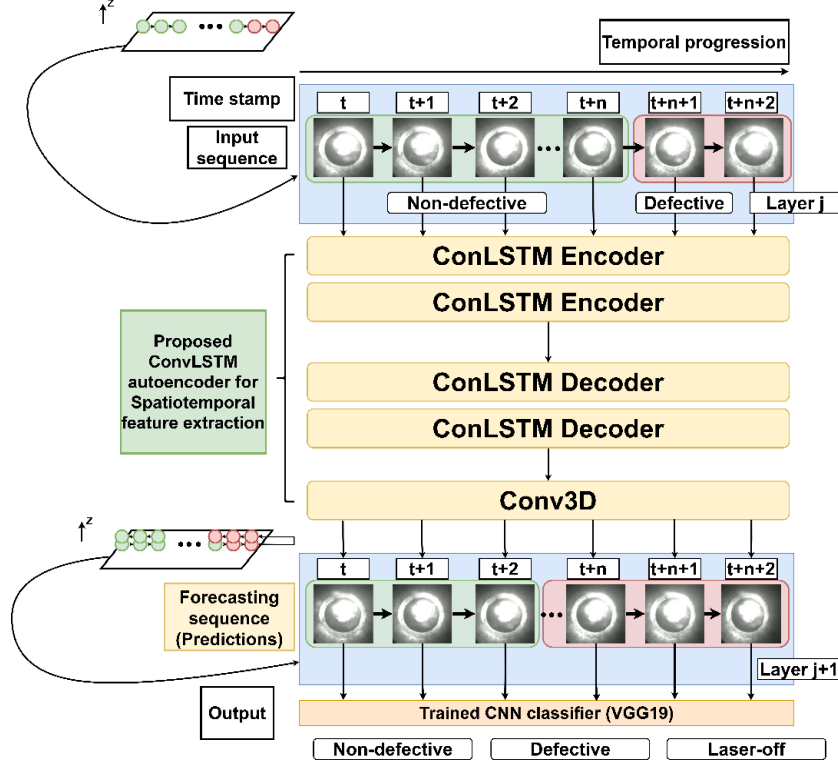


FIGURE 4: THE WORKFLOW OF PROPOSED CONVLSTM-CNN ARCHITECTURE FOR DEFECTS PREDICTION IN SINGLE-BEAD WALL PRINTING, WITH INPUT SEQUENCE, FORECASTING SEQUENCE, TIME STAMPS, AND FINAL OUTPUT

For the development of the ConvLSTM-CNN model in this study, co-axial melt pool images from the single-bead wall experiments are utilized. Both the melt pool images, and the corresponding time stamps are sampled from the recorded videos at a frequency of 10 Hz, resulting in one image every 100 milliseconds. Each image was labelled with its time stamp to encode its temporal information. It is worth noting that in every layer of printing, there are preprogrammed laser-off periods at the beginning and the end of each layer to improve process stability.

TABLE 1: GEOMETRIES OF PRINTED COMPONENTS, PROCESS PARAMETERS

Parameters	Values
Geometry	Single-bead wall
Dimension	90 mm * 42.5 mm
Number of layers per sample	50
Laser power (kW)	[2.3,2.53]
Laser beam diameter (mm)	2
Speed (mm/s)	[25,27.5]
Powder flow rate (g/min)	12
Dwell time between layers (s)	[0.5,1.0]
Layer thickness (mm)	0.85
Stand-off distance (mm)	12
Types of defects generated	Cracks, keyhole pores

Therefore, the laser-off periods are used to segment the image data into layers. Within each layer, the images are sequenced chronologically according to the time stamps. Additionally, every image sequence is organized into fixed-length segments of 10 images. All images are resized to 32×32 pixels, ensuring the consistency for effective ConvLSTM autoencoder training and evaluation.

2.5 Model Development

The proposed sequence-to-sequence ConvLSTM model is designed to extract spatiotemporal correlations from past (input) observations and generate future trajectories. The architecture comprises two ConvLSTM encoder layers, two ConvLSTM decoder layers, and a three-dimensional convolutional layer (Conv3D). These three components compose the ConvLSTM autoencoder, which is utilized to capture spatiotemporal features. The results of the ConvLSTM autoencoder and visualization of the forecasting melt pool images are presented in Section 3. The detailed architecture of ConvLSTM-CNN model is shown in FIGURE 4. The data used is the melt pool images labeled with time stamps and sequenced according to the time stamp. Every image in sequence is fed into the ConvLSTM cell in the encoder layer to extract the spatiotemporal features and followed by the decoder layer and Conv3D layer to generate the forecasting sequence of next-layer melt pool images. The forecasting sequence of melt pool images represents the predicted next-layer printing state. Every image in the forecasting sequence will be fed into a pre-trained CNN classifier. In this study, a pre-trained

Visual Geometry Group (VGG) 19 CNN network is used for the melt pool image classification. As a result, the defect prediction is achieved by giving the classification of every image in the forecasting sequence. For the next-layer prediction task, the ConvLSTM successfully captured the spatial influences from previous layer.

3. RESULTS AND DISCUSSION

In this section, the results of the preliminary study using the ConvLSTM model are presented and discussed. The results are based on single-bead wall printing experiments. In this study, the training loss is the deviation of pixel-wise mean square error (MSE). The model training hyperparameters and hardware information are listed in TABLE 2.

TABLE 2: MODEL TRAINING HYPERPARAMETERS AND HARDWARE INFORMATION

Parameters	Values
Learning rate	0.0001
Decay rate 1	0.9
Decay rate 2	0.98
Weight initialization	zeros
Operating system	Ubuntu 20.02
Computer Processing Unit (CPU)	i7-8750H
CPU Speed	2.20 GHz
On-board Random Access Memory (RAM)	2 X 8 GB
Graphic Processing Unit (GPU)	GTX1060
GPU dedicated RAM	8 GB

The forecasting image sequence generated by the ConvLSTM autoencoder is shown in FIGURE 5, which provides a visualization of the model's forecasting performance. Specifically, the image grids in FIGURE 5 have fixed columns, representing the consistent time length for predictions. FIGURE 5 consists of two parts: input sequence and prediction, ground truth, and difference grids. The grid part consists of three rows: (1) the model's prediction sequence, (2) the ground-truth sequence, and (3) a pixel-wise difference between first row and second row highlighting discrepancies between predictions and actual observations. From the difference row, the predictions are slightly different from the ground truth after 1000 training steps, which emphasizes the performance of the ConvLSTM autoencoder.

FIGURE 6 presents the training loss and the validation loss and accuracy curves. The curves demonstrate the convergence behaviour of the ConvLSTM autoencoder during training. Initially, both the training and validation losses decrease rapidly, indicating effective learning of spatiotemporal patterns from the input sequences. However, after several epochs, the validation loss plateaus, while the training loss continues to decline,

suggesting potential overfitting. Despite this, the validation accuracy remains relatively stable, indicating that the model can generalize to unseen sequences to some extent. Further regularization techniques, such as dropout or early stopping, may improve generalization performance.

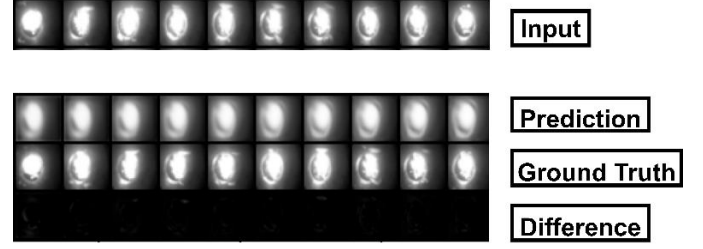


FIGURE 5: INPUT SEQUENCE OF IMAGES, GENERATED PREDICTION SEQUENCE, GROUND TRUTH, AND DIFFERENCE OF PREDICTION AND GROUND TRUTH AFTER 1000 TRAINING STEPS.

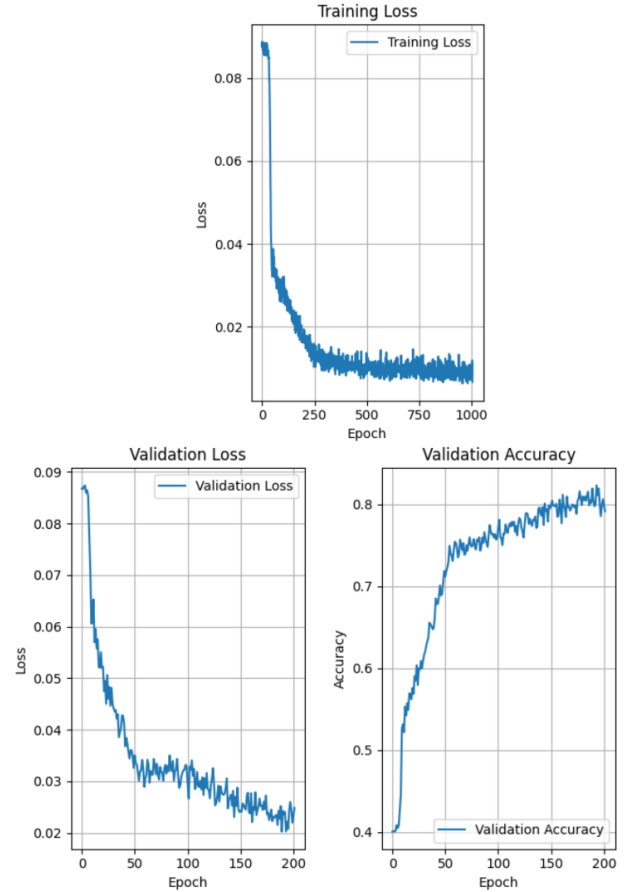


FIGURE 6: TRAINING LOSS, VALIDATION LOSS AND VALIDATION ACCURACY OF CONVLSTM AUTOENCODER.

To achieve the defect prediction, the predicted melt pool image sequence of the ConvLSTM autoencoder is further fed into a fine-tuned classifier model. In this study, the next-layer melt pool images are classified by the VGG19 CNN for non-defective, defective, and laser-off classifications. The fine-tuned

VGG19 CNN classifier confusion matrix is shown in FIGURE 7. The least True Positive (TP) percentage is over 90%.

To intuitively view the work of the classifier, Figure 8 compares two images of non-defective and defective melt pools. Unlike non-defective cases, defective melt pools exhibit irregularities in size and shape, deviating from a circular morphology. In the figure, the melt pool contour is highlighted in red solid line, while blue dashed and solid lines indicate smoke and spatter formation, respectively. Although melt pool shapes are inherently irregular due to the dynamic nature of powder-based deposition, another feature of defective melt pools is its size tends to be larger. Furthermore, material evaporation, smoke, and spatters are more significant in defective melt pool images than in non-defective ones (blue lines). In addition to these visible differences, subtle visual features, though less intuitive, also contribute to the VGG19 CNN classifier's ability to distinguish between defective and non-defective melt pools.

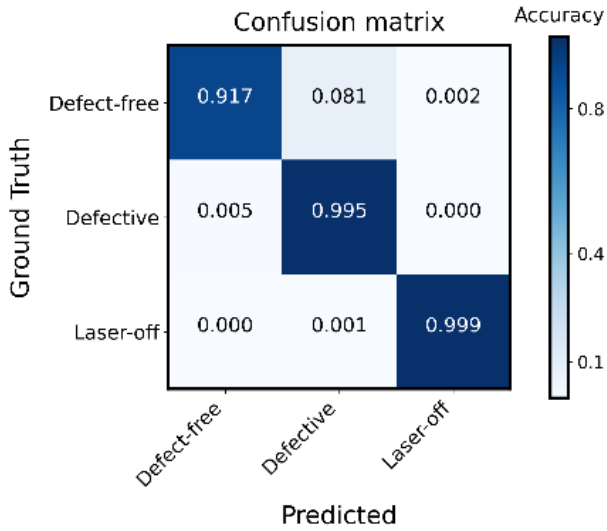


FIGURE 7: CONFUSION MATRIX OF THE PRE-TRAINED BINARY CLASSIFIER FOR THE MELT POOL IMAGES.

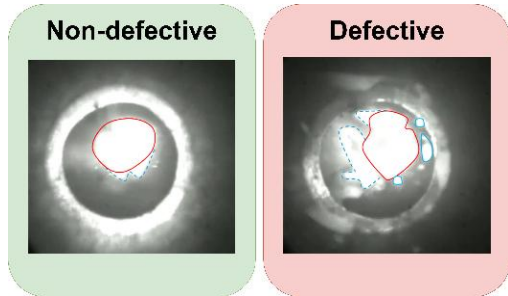


FIGURE 8: COMPARISON BETWEEN TWO SAMPLES OF MELT POOL IMAGES FROM NON-DEFECTIVE AND DEFECTIVE PROCESSES (RED LINE: MELT POOL CONTOUR, BLUE LINE: EVAPORATIONS AND SPATTERS).

Overall, the results from the preliminary study highlight both the strengths and limitations of the ConvLSTM autoencoder

in melt pool image predictions. Although the model successfully captures temporal dependencies and produces reasonable next-layer image sequences in single-bead wall printing, it struggles with complex spatial relationships, particularly for non-linear printing patterns in rectangular and square geometries due to the model's exclusion of the zig-zag printing pattern, and the complex spatial influences inherent in the L-DED process. In these geometries, multiple adjacent tracks from both previous and current layers influence the printing outcomes, as shown in FIGURE 3. The ConvLSTM model, which contains limited contextual window, struggles to account for these intricate spatial interactions and relatively larger temporal information. To address this limitation, an approach that integrates comprehensive spatiotemporal relationship factors within the LSTM model is essential to improve the defect prediction accuracy, which would be the future work of this study.

4. CONCLUSION

In this study, the key contribution is the novel framework of ConvLSTM-CNN model for forecasting the next-layer defect occurrence in the L-DED process. The preliminary validation experiment from single-bead wall printing demonstrated the ConvLSTM autoencoder's capability to capture the spatiotemporal relationships across layers and forecast next-layer melt pool image sequences. The novelty lies within the proposed ConvLSTM-CNN model shifting the reactive defect detection to proactive defect prediction by generating and classifying next-layer melt pool images in L-DED. In conclusion, this work lays the foundation for adaptive process adjustments and preemptive defect mitigation in L-DED, contributing to adaptive and intelligent process control in advanced manufacturing. Future work will explore the solutions for predicting more complex geometries such as multi-layer multi-track printing.

ACKNOWLEDGEMENTS

This work is supported by the Agency for Science, Technology, and Research (A*STAR), and National University of Singapore (NUS). It is also supported by A*STAR Graduate Scholarship (AGS) funded by A*STAR.

REFERENCES

- [1] L. Chen *et al.*, "In-situ process monitoring and adaptive quality enhancement in laser additive manufacturing: A critical review," *J. Manuf. Syst.*, vol. 74, pp. 527–574, Jun. 2024, doi: 10.1016/j.jmsy.2024.04.013.
- [2] A. Ben Hammouda, H. Mrad, H. Marouani, A. Frikha, and T. Belem, "Process Optimization and Distortion Prediction in Directed Energy Deposition," *J. Manuf. Mater. Process.*, vol. 8, no. 3, p. 116, May 2024, doi: 10.3390/jmmp8030116.
- [3] J. Mi *et al.*, "In-situ monitoring laser based directed energy deposition process with deep convolutional neural network," *J. Intell. Manuf.*, vol. 34, no. 2, pp. 683–693, Feb. 2023, doi: 10.1007/s10845-021-01820-0.
- [4] M. Khanzadeh, S. Chowdhury, M. A. Tschopp, H. R. Doude, M. Marufuzzaman, and L. Bian, "In-situ monitoring

- of melt pool images for porosity prediction in directed energy deposition processes,” *IISE Trans.*, vol. 51, no. 5, pp. 437–455, May 2019, doi: 10.1080/24725854.2017.1417656.
- [5] M. Atwya and G. Panoutsos, “In-situ porosity prediction in metal powder bed fusion additive manufacturing using spectral emissions: a prior-guided machine learning approach,” *J. Intell. Manuf.*, vol. 35, no. 6, pp. 2719–2742, Aug. 2024, doi: 10.1007/s10845-023-02170-9.
- [6] Z. Tang *et al.*, “Investigation on coaxial visual characteristics of molten pool in laser-based directed energy deposition of AISI 316L steel,” *J. Mater. Process. Technol.*, vol. 290, p. 116996, Apr. 2021, doi: 10.1016/j.jmatprotec.2020.116996.
- [7] H. Wang, B. Li, and F.-Z. Xuan, “Acoustic emission for in situ process monitoring of selective laser melting additive manufacturing based on machine learning and improved variational modal decomposition,” *Int. J. Adv. Manuf. Technol.*, vol. 122, no. 5–6, pp. 2277–2292, Sep. 2022, doi: 10.1007/s00170-022-10032-6.
- [8] T. Hauser, R. T. Reisch, T. Kamps, A. F. H. Kaplan, and J. Volpp, “Acoustic emissions in directed energy deposition processes,” *Int. J. Adv. Manuf. Technol.*, vol. 119, no. 5–6, pp. 3517–3532, Mar. 2022, doi: 10.1007/s00170-021-08598-8.
- [9] L. Chen *et al.*, “In-situ crack and keyhole pore detection in laser directed energy deposition through acoustic signal and deep learning,” *Addit. Manuf.*, vol. 69, p. 103547, May 2023, doi: 10.1016/j.addma.2023.103547.
- [10] L. Chen and S. K. Moon, “In-situ defect detection in laser-directed energy deposition with machine learning and multi-sensor fusion,” *J. Mech. Sci. Technol.*, vol. 38, no. 9, pp. 4477–4484, Sep. 2024, doi: 10.1007/s12206-024-2401-1.
- [11] L. E. Criaes, Y. M. Arisoy, B. Lane, S. Moylan, A. Donmez, and T. Özel, “Laser powder bed fusion of nickel alloy 625: Experimental investigations of effects of process parameters on melt pool size and shape with spatter analysis,” *Int. J. Mach. Tools Manuf.*, vol. 121, pp. 22–36, Oct. 2017, doi: 10.1016/j.ijmachtools.2017.03.004.
- [12] J. Kim, Z. Yang, H. Ko, H. Cho, and Y. Lu, “Deep learning-based data registration of melt-pool-monitoring images for laser powder bed fusion additive manufacturing,” *J. Manuf. Syst.*, vol. 68, pp. 117–129, Jun. 2023, doi: 10.1016/j.jmsy.2023.03.006.
- [13] D. Y. Kononenko, V. Nikonova, M. Seleznev, J. Van Den Brink, and D. Chernyavsky, “An in situ crack detection approach in additive manufacturing based on acoustic emission and machine learning,” *Addit. Manuf. Lett.*, vol. 5, p. 100130, Apr. 2023, doi: 10.1016/j.addlet.2023.100130.
- [14] S. M. Estalaki, C. S. Lough, R. G. Landers, E. C. Kinzel, and T. Luo, “Predicting defects in laser powder bed fusion using in-situ thermal imaging data and machine learning,” *Addit. Manuf.*, vol. 58, p. 103008, Oct. 2022, doi: 10.1016/j.addma.2022.103008.
- [15] L. Chen *et al.*, “Multisensor fusion-based digital twin for localized quality prediction in robotic laser-directed energy deposition,” *Robot. Comput.-Integr. Manuf.*, vol. 84, p. 102581, Dec. 2023, doi: 10.1016/j.rcim.2023.102581.
- [16] B. Abranovic, S. Sarkar, E. Chang-Davidson, and J. Beuth, “Melt pool level flaw detection in laser hot wire directed energy deposition using a convolutional long short-term memory autoencoder,” *Addit. Manuf.*, vol. 79, p. 103843, Jan. 2024, doi: 10.1016/j.addma.2023.103843.
- [17] H. Ko, J. Kim, Y. Lu, D. Shin, Z. Yang, and Y. Oh, “Spatial-Temporal Modeling Using Deep Learning for Real-Time Monitoring of Additive Manufacturing,” in *Volume 2: 42nd Computers and Information in Engineering Conference (CIE)*, St. Louis, Missouri, USA: American Society of Mechanical Engineers, Aug. 2022, p. V002T02A019. doi: 10.1115/DETC2022-91021.
- [18] S. Yu *et al.*, “Thermo-fluid dynamics and morphology evolution of nickel-based superalloy in a multilayer laser directed energy deposition process,” *J. Mater. Res. Technol.*, vol. 28, pp. 1276–1293, Jan. 2024, doi: 10.1016/j.jmrt.2023.12.047.
- [19] V. Errico, S. Campanelli, A. Angelastro, M. Dassisti, M. Mazzarisi, and C. Bonserio, “Coaxial Monitoring of AISI 316L Thin Walls Fabricated by Direct Metal Laser Deposition,” *Materials*, vol. 14, no. 3, p. 673, Feb. 2021, doi: 10.3390/ma14030673.