



Partial-physics-informed multi-fidelity modeling of manufacturing processes

Jeremy Cleeman¹, Kian Agrawala¹, Evan Nastarowicz, Rajiv Malhotra^{*}

Department of Mechanical and Aerospace Engineering, Rutgers University, Piscataway, NJ 08854, USA

ARTICLE INFO

Associate Editor: Marion Merklein

Keywords:

Process modeling
Physics-Informed Machine Learning
Multi-fidelity Learning
Transfer Learning
Training data cost

ABSTRACT

The design and control of manufacturing processes hinges on predictive modeling of its parametric effects. The deployability of Machine Learning (ML) models has made them of increasing interest for this purpose. Recent work has addressed the experimental and computational costs of creating the requisite training data. But these methods incur a physics development cost due to the need to intuitively and iteratively derive accurate physics-based process models for generating the training data. This issue is rarely addressed in the literature. This paper describes a Science-informed multifidelity-aided reduced-cost Machine Learning (Smart-ML) approach to tackle this challenge. The novelty lies in relaxing the existing constraint that the physics-based source in multi-fidelity learning must qualitatively match the experimental ground truth while constraining the source to use conservation laws. An additive, a subtractive, and a hybrid additive-deformative process with different levels of physical understanding are used as testbeds to demonstrate that Smart-ML can reduce the physics development cost by multiple human-years, experimental cost by as much as 60 %, and computational cost by orders of magnitude. These results are discussed in the context of how the proposed approach can move ML beyond the creation of copies of known physics-based models towards the accelerated and inexpensive derivation of functionally new ML-based process models from partially known physics and small experimental datasets.

1. Introduction

Innovative manufacturing processes can have a disruptive socio-economic impact by going beyond existing limitations on cost, flexibility, throughput, sustainability, and product functionality. Past examples include the reduction in automobile costs via stamping in the 19th century (Hounshell, 1984), the achievement of complex geometries for hard-to-form sheet metal via superplastic forming in the 20th century (Barnes, 2007), and the disruption of batch size-cost tradeoffs by additive manufacturing in the 21st century (Huang et al., 2013). More recent examples include advances in scalable micro/nanoscale additive manufacturing of polymers based on massive multiplexing of light (Saha et al., 2019); metamorphic approaches to flexible manufacturing that obviate part-shape-specific tooling (Daehn and Taub, 2018); hybrid processes that combine additive manufacturing and sheet/bulk metal forming (Merklein et al., 2021); nanoscale deformation of 2D materials to control their electronic behavior (Yi et al., 2022); approaches that combine material deposition, forming and sintering to create surface-conformal electronics (Devaraj and Malhotra, 2019); and

dynamic deposition for massively multiplexed additive manufacturing (Cleeman et al., 2022).

Such innovations often utilize new physical phenomena to realize their technological advantage which, in turn, impose unique and complex parametric effects on the part's geometric and property attributes. Predictive modeling of such parametric effects is crucial for design and control of manufacturing processes and systems. Machine Learning (ML) models are of increasing interest for this purpose since they enable precise real-time predictions for complex multivariate parametric interactions and are easy to deploy industrially (Arinez et al., 2020). This is a significant advantage over spatiotemporally discretized physics-based computational process models which have limited real-time predictive capability. ML also has an edge over more computationally efficient physics-based analytical models since such models are not always derivable. Even when they can be derived they are often only partially complete since they cannot accurately capture the entire spectrum of parametric effects in many cases.

Using experiments to generate the data needed to train the ML models incurs a cost C_{EXP} proportional to the time, material, and human

^{*} Corresponding author.

E-mail address: rajiv.malhotra@rutgers.edu (R. Malhotra).

¹ Authors made an equal contribution

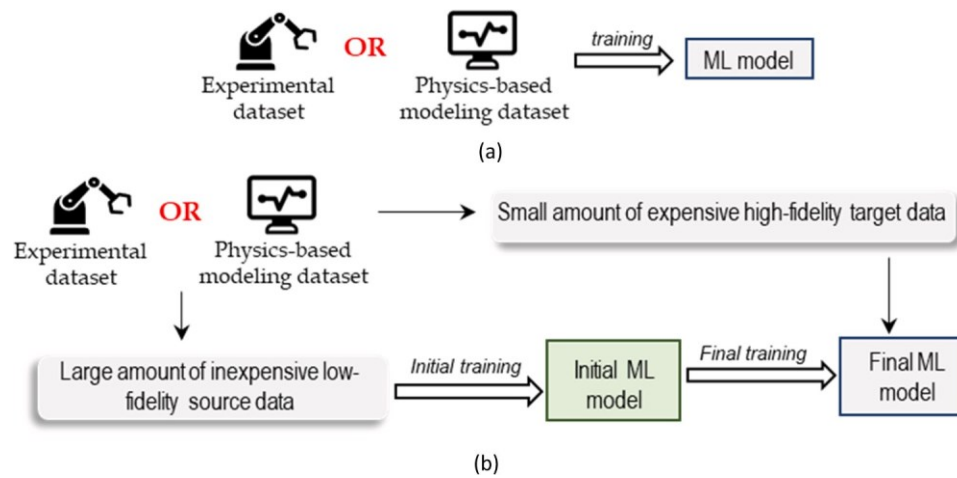


Fig. 1. Schematic of (a) Direct learning (b) Multi-fidelity learning.

resources needed for experimentation. The computational cost C_{COMP} , in CPU-hours, is the computational effort needed to perform physics-based process simulations for generating the training data. But there is another, often-ignored, physics development cost C_{DEV} which consists of the time and resources needed for iterative and intuitive human derivation of accurate physics-based process models from which data is extracted to train ML models. This time and resources are necessary because the fidelity of the training data dictates the ability of the trained ML model to reflect reality. Thus, physics-based models that are used as the source of training data must have sufficient qualitative and quantitative agreement with the experimental ground truth. The quantitative calibration of physics model parameters has a small contribution to C_{DEV} since multiple methods are available to perform this model calibration in an automated and rapid fashion. The primary component of C_{DEV} is the time and human resources expended in the manual and iterative identification of the physical constituents, the functional form of the multiphysical and multiscale linkages, and the mathematical form of the constitutive laws that constitute physics-based process models with sufficient accuracy to supply training data for ML.

Since the physics-based models needed to generate training data for ML models are often derived by humans in an iterative and intuitive fashion this paper uses human-years as a metric of C_{DEV} . The C_{DEV} is often on the order of many human-years of effort for new processes. Consider incremental sheet metal forming. This process saw renewed interest in academia in 2005 (Dufloy et al., 2018) but it took till the mid-2010s to create accurate physics-based predictive models of in-process metal fracture and models of formed material properties like fatigue are still unavailable. Fused Filament Fabrication has seen wide usage since 2000 but models for predicting the printed road width have taken nearly two more decades, till 2019 for analytical models (Agassant et al., 2019) and till 2018 for computational fluid dynamics models (Serdeczny et al., 2018). Different metal additive manufacturing processes can yield different hardness and fatigue properties for the same material and the physics modeling of the underlying parametric effects is still an active area of research, even though these processes were first developed more than a decade ago (Mukherjee and DebRoy, 2019). The sintering of printed nanowires for highly conductive printed electronics was reported as early as 2009 (Fan et al., 2009). But the development of physics-based models that can predict the mesoscale electrical conductivity in terms of the sintering parameters has taken till 2021 (Devaraj et al., 2021).

The relative magnitude of C_{EXP} , C_{COMP} , and C_{DEV} depends on the approach used for training the ML model. The first training method, called direct learning here, performs training using a single dataset from a single fount of information (Fig. 1a). Direct learning can be performed on purely experimental data, numerous examples of which can be found

in a recent review by (Arinez et al., 2020). This results in high C_{EXP} despite the development of sparse sampling methods and niche cases of high-throughput experimental testing, e.g., for measuring parametric effects on the part's microstructure in metal additive manufacturing (Pegues et al., 2021).

This experimental approach is increasingly being replaced by direct learning on data generated from physics-based models, which reduces C_{EXP} to the experimental effort needed for calibration and validation of the physics-based models (Zhu et al., 2021). trained Physics-informed Neural Networks (PINNs) on an experimentally calibrated high-fidelity finite element model to predict the temperature dynamics of the melt pool in metal additive manufacturing. The use of PINNs reduced C_{COMP} by decreasing the number of labeled model-generated samples needed for training and minimized C_{EXP} since experiments were only needed for calibrating and validating the finite element model. While this model can be used to study the effects of the process parameters on the weld pool size, the availability of sufficient physics and accurate constitutive thermal laws for embedding in the PINNs was contingent on the expenditure of sufficient C_{DEV} for model development in the author's previous work (Yan et al., 2018). Similarly, a Graph Neural Network approach for modeling the mesoscale temperature history during Directed-Energy-Deposition was trained on a finite element model by (Mozaffar et al., 2021). This incurred a C_{DEV} related to the identification of the appropriate constitutive laws for the underlying heat and mass transfer equations in (Mozaffar et al., 2019). Microstructure evolution during metal additive manufacturing was modeled using a physics-embedded graph network by (Xue et al., 2022), building on the significant C_{DEV} incurred over the years in both phase-field modeling and metal additive manufacturing. The employment of analytical models to generate the training data reduces the C_{COMP} relative to the use of computational models. For example, (Kapusuzoglu and Mahadevan, 2020) trained different ML models on data from a new analytical model for accurately predicting the porosity and the bond strength in Fused Filament Fabrication. This analytical model was built on past efforts to incorporate actual filament geometry and the changes in it induced by printing. This incurred a C_{DEV} in addition to that from past literature, which stretched back 23 years (Pokluda et al., 1997), in order to achieve the requisite accuracy on which the ML model could be trained. Further, the derivation of analytical models is not always possible for a given process or for a given geometric or property metric. Overall, direct learning on model-generated data incurs a high C_{DEV} that can significantly overshadow the more commonly addressed C_{EXP} and C_{COMP} .

The second training approach is called multi-fidelity learning (Fig. 1b). It trains a ML model on a large, inexpensive, but inaccurate source dataset to satisfy the need for large amounts of data and then fine-

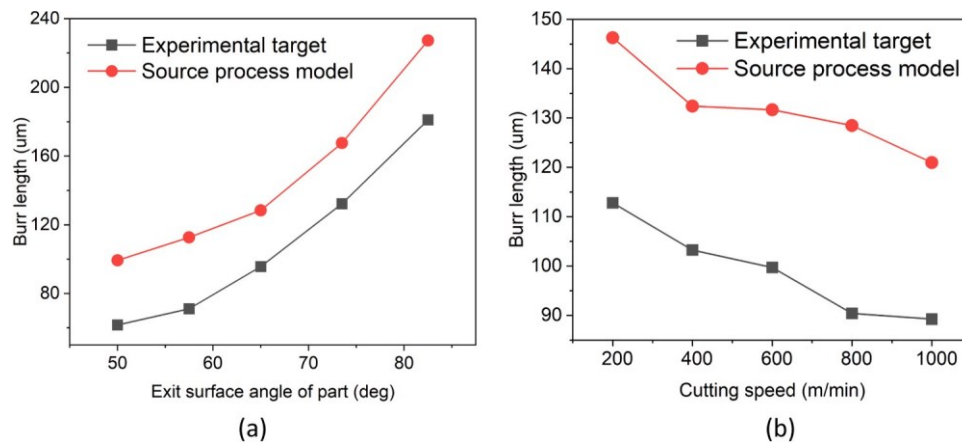


Fig. 2. Comparison between experimental target and source process model for the effect of different process parameters on burr length in milling (Liu et al., 2022).

tunes it on a smaller, more expensive, and more accurate target dataset to capture the ground truth (Peherstorfer et al., 2018). Computational or analytical process models can be used as the target with process models of lower fidelity as the source. (Ramezankhani et al., 2021) generated the source and target data for a composites curing process using previously developed Finite Element Models. (Menon et al., 2022) used a similar approach for modeling the melt pool size in laser-direct energy deposition using Gaussian Processes. An Eager-Tsai model, based on solution of the 3D temperature distribution produced by a traveling distributed heat source moving on a semi-infinite plate, was used as the source. A non-linear decoupled 3D transient FEM solver with Marangoni convection, convection and radiation heat losses, and the temperature-dependent thermophysical properties omitted by the Eager-Tsai model, was used as the target. (Huang et al., 2021) used a similar approach for melt pool prediction in electron beam additive manufacturing of metals with Rosenthal's model as the source and a Finite Element Model as the target, but with a point neural network as the ML model. A key observation is that these multi-fidelity learning approaches are predicated on high C_{DEV} for a new process since they are contingent on the availability of a target process model of qualitative completeness and quantitative accuracy sufficient to represent the experimental ground truth.

A physics-based process model and experiments can be used as source and target respectively. This ensures that the ground truth, i.e., experimental observations, are always captured by the final ML model while reducing C_{EXP} by decreasing the number of experimental samples needed. (Alam et al., 2020) performed optimization of Fused Filament Fabrication parameters for printing of log-pile metamaterials by using a Gaussian process as the ML model, a Finite Element Mechanics model as the source, and experimental measurements of the printed part's mechanical response as the target. Multi-fidelity learning of a feedforward neural network was performed to reduce the experimental cost of creating models that predicted the part weight as a function of the injection molding parameters across different materials (Lockner et al., 2022) and different part sizes and machine types (Lockner and Hopmann, 2021). The source was a computational fluid dynamics model with pre-identified constitutive laws and numerical techniques. (Liu et al., 2022) used Finite Element simulations of burr formation during milling as the source to create a 1D residual convolutional neural network and corrected this learnt ML model to match experimental data using transfer learning. (Wang et al., 2021) used a finite element model as the source for predicting milling forces with a feedforward neural network. (Saunders et al., 2022; Saunders et al., 2023) used multiple sources, in the form of analytical and computational models of different fidelities, with a Gaussian Process ML model of the melt pool size and microstructure in laser powder bed fusion of metals. Other approaches have used analytical models, when available, as the source, e.g., the

prediction of surface roughness in CNC machining (Misaka et al., 2020).

Despite its advantages, there is a fundamental issue that plagues this last multi-fidelity learning approach. It is the fact that the physics-based process models used as the source are developed to a point that ensures a qualitative match between the source and the dataset. Thus, transfer learning is only used to perform quantitative corrections such as scaling and translation in the output space. This requires human identification of the appropriate physical constituents of the process models, of adequate constitutive laws to link these constituents, and of appropriate numerical techniques for performing the simulations. Since the requisite physical knowledge is often lacking for new manufacturing processes this multi-fidelity learning approach is unable to reduce C_{DEV} despite having the ability to decrease C_{EXP} . Fig. 2 shows an example of this by-design similarity between the source model and the experimental target from past work on burr prediction in milling (Liu et al., 2022). The C_{DEV} in this case is incurred because the Finite Element Model of material removal that predicts the necessary training data for the ML model requires derivation of constitutive laws for strain hardening, strain-rate sensitivity, and thermal softening, the form of the friction model, and the iterative identification of appropriate mesh refinement and element choices.

Overall, the issue of high C_{DEV} for new processes is not addressed by the state-of-the-art. The question is, is it possible to derive accurate ML-based process models of parametric relationships in manufacturing processes while simultaneously reducing the need for highly accurate and functionally complete process physics models to create the training data (thus reducing C_{DEV}), reducing the amount of experimentally generated training data (thus reducing C_{EXP}), and reducing the computational effort needed to create the training data (thus reducing C_{COMP}). This paper addresses this question by modifying the second multi-fidelity learning method into an approach called Science-informed multifidelity-aided reduced-cost Machine Learning (Smart-ML). The novelty is to relax the constraint that a qualitative match between the source and target data is necessary and to correct the resulting discrepancy using transfer learning. The magnitudes of the different training costs and prediction accuracy for Smart-ML are compared to that for direct learning on only experimental data for three distinct types of manufacturing processes whose working principles span material addition, removal, and deformation.

Using these processes for demonstration purposes also tackles another issue. The larger vision for Smart-ML is that it will become the go-to approach for training ML models of parametric effects without prior calculation of the reduction in C_{DEV} . This is because such a computation requires the creation of an accurate physics-based model and therefore the expenditure of the very C_{DEV} that Smart-ML aims to save. In fact, the projected savings of C_{EXP} and C_{COMP} in any form of multifidelity learning including Smart-ML are subject to a similar

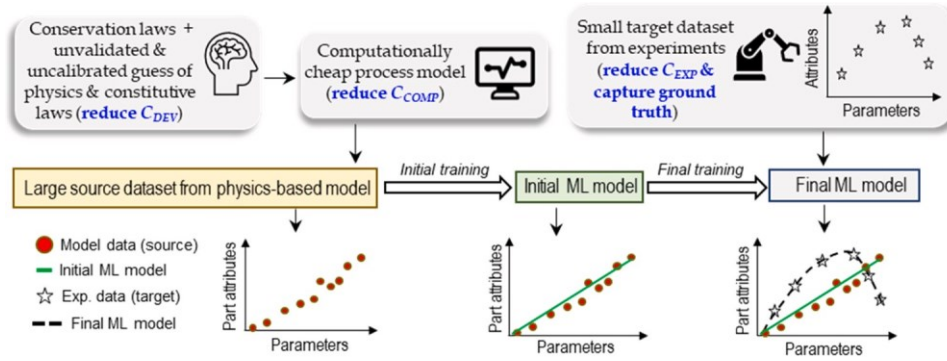


Fig. 3. Proposed Smart-ML approach.

critique since closed-form a-priori calculation of the cost savings in multifidelity learning is fundamentally impossible without sacrificing the savings themselves. Thus, the adoption of multifidelity learning in the literature is driven by empirical proof of cost savings across different use cases, e.g., in image and signal processing for manufacturing processes with Convolutional Neural Networks as the base machine learning model. This indicates that realizing our vision for Smart-ML requires empirical proof of the reduction in C_{DEV} , C_{EXP} and C_{COMP} via Smart-ML. This work performs this examination by quantifying these costs as described in Section 2.2 for each of the three process testbeds that span different physical principles and different levels of physical understanding.

The following sections discuss the details of the Smart-ML approach, followed by a description of the source process models and the experimental methods, and ending with the results and a discussion of the larger impact of Smart-ML on manufacturing processes.

2. Methods

2.1. Machine learning techniques

2.1.1. Overview of Smart-ML

Fig. 3 illustrates the Smart-ML method. The physics-based process models that constitute the source are based on the following principles. First, the source process model should include one or more conservation equations to retain conformance to the basic principles of nature. This is the science-informed component of Smart-ML that regularizes the ML model within the constraints of existing knowledge of the process physics. Second, the source process model should use an initial guess for the physical phenomena to be included and for the constitutive laws to be used without any trial-and-error based qualitative or quantitative calibration against experiments. This is a key departure from the existing literature on multi-fidelity learning in manufacturing. It obviates the need for the source and the target to be qualitatively similar, thus eliminating the trial-and-error construction of the process model, with the goal of reducing C_{DEV} for new processes for which the knowledge of the appropriate physics and constitutive laws is lacking. Finally, the degree of spatiotemporal discretization used in the source model should not be so high that it causes the C_{COMP} to exceed the user's budget. This budget is typically a finite non-zero value that is based on the resources and time available for computational generation of training data. If possible, purely analytical models without spatiotemporal discretization should be used as the source since such models often need much lesser computational effort than discretization-based models like Finite Element Analysis. The ML model trained on the source is then updated on a much smaller experimental dataset (i.e., target) than that necessary for direct learning with only experimental data by using transfer learning. This reduces C_{EXP} as compared to direct learning while ensuring that the final ML model's predictions accurately capture the experimental ground truth.

2.1.2. Base machine learning model

This paper uses epsilon-Support Vector Regression (ϵ -SVR) as the base ML model on which transfer learning is performed. Other regression ML approaches can be used without loss in generality. The mechanics of the ϵ -SVR is described briefly here and the reader is referred to past theoretical developments for further details (Drucker et al., 1996). An ϵ -SVR aims to find the function that maps the inputs to the outputs within an error band ϵ . If x and $f(x)$ denote the inputs and outputs respectively and ω are the weights that perform the mapping then the mapping function can be written as in Eq. (1). The constant b is solved using the Karush-Kuhn-Tucker conditions. The vector of weights ω , and the slack variables ξ_i and ξ_i^* used to cope with the fact some deviations from the margin of error ϵ might have to be tolerated, are computed by minimizing the expression in Eq. (2). The value of C determines the penalty applied to violation of the error band. The Gaussian Radial Basis Function (RBF) was used here for the kernel trick, thus allowing the linear ϵ -SVR model to capture nonlinear functional relationships (Jain et al., 2014). In this work the values of C and ϵ for the ϵ -SVR were based on brute force identification of parameter combinations that maximized the model performance for a given source dataset and were not changed during transfer learning.

$$f(x) = \langle \omega, x \rangle + b \quad (1)$$

$$\frac{1}{2} \|\omega\|^2 + C \sum_{i=1}^l (\xi_i + \xi_i^*) \quad (2)$$

2.1.3. Transfer Learning approach

Transfer learning updates the ML model based on the difference between the source and the target. Consider a domain $D = (X, P_X)$ with features X and a marginal probability distribution P_X , and a task $T = (Y, f(\cdot))$ that consists of a label space Y (continuous values here) and a function $f(\cdot)$ which is trained to predict Y . For a source domain D_s and learning task T_s and a target domain D_t and learning task T_t , where $D_s \neq D_t$ or $T_s \neq T_t$, transfer learning reduces the amount of data needed to learn $f_t(\cdot)$ by using the a-priori trained $f_s(\cdot)$. This can be performed by reweighting the training data (instance-based transfer learning), altering the ML model's weights (parameter-based transfer learning), or leveraging common features between the source and target (feature-based transfer learning). The reader is referred to the excellent review article by (Pan and Yang, 2010) for a deeper survey of the different kinds of transfer learning.

The instance-based transfer learning method used here is TrAdaBoost.R2 by (Pardoe and Stone, 2010) since it is suitable for regression tasks. This method derives from the AdaBoost family of techniques in which each target and source instance receives a weight used for training. The weight indicates the relative importance of each target or source instance. The instances are reweighted after each training iteration, with target instances that are incorrectly predicted by the ML model trained in the previous iteration receiving larger weights. Thus,

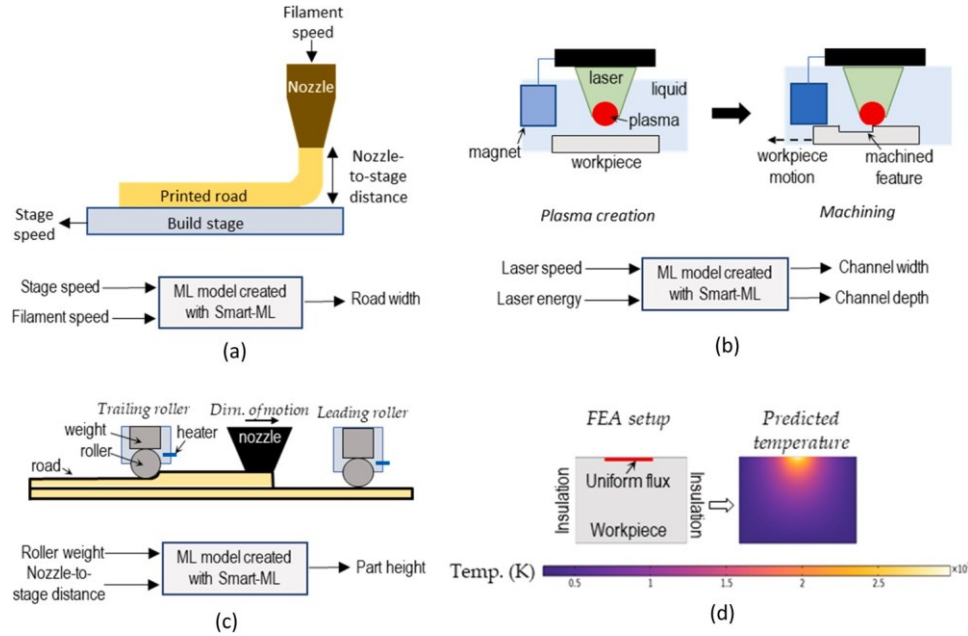


Fig. 4. Schematic of the process and associated ML model for the (a) FFF problem (b) M-LIPMM problem (c) In-situ rolling problem. (d) Example of peak temperature distribution for one laser pulse predicted by the M-LIPPM source.

learning iteratively focuses on those instances that are most difficult to predict correctly. The TrAdaBoost.R2 variant of AdaBoost takes the target and source datasets and combines them into a single training dataset. Further, in each boosting iteration TrAdaBoost increases the weights of target instances that are misclassified while reducing the weights of the corresponding source instances.

The basic mechanics of this method is described here and the reader is referred to (Pardoe and Stone, 2010) for greater depth. Each boosting iteration consists of the following steps and the iterations are repeated till a user-specified number of iterations N is reached. The weights of the source and target instances w_S and w_T are normalized and the normalized error vectors for the source and the target ε_S and ε_T are computed. The total weighted error for the target dataset E_T is computed as in Eq. (3), where n_T is the number of target instances. The source weights are reduced and the target weights increased based on the error, calculated as in Eqs. 4 and 5, where n_S is the number of source instances. Here, the number of boosting iterations N was fixed at thirty.

$$E_T = \frac{1}{n_T} w_T^T \varepsilon_T \quad (3)$$

$$w_S = w_S \beta_S^{\varepsilon_S} \text{ and } w_T = w_T \beta_T^{-\varepsilon_T} \quad (4)$$

$$\beta_S = \frac{1}{\sqrt{s/l_i} + 2 \ln(n)} \text{ and } \beta_T = E_T (1 - E_T) \quad (5)$$

2.1.4. Training and testing

The following approach was used for training and testing and to examine the advantages of Smart-ML relative to direct learning. First, direct learning of the ε -SVR was performed on only the experimental target data in an incremental fashion. In each iteration a progressively greater number of training samples were used for training till the Root Mean Square Error (RMSE) on the testing data, i.e., withheld portion of the experimental dataset, did not decrease further. This was done 1000 times using random sampling for obtaining training and testing datasets in a 90:10 ratio, thus yielding the mean and standard deviation of the smallest error δ_d and the corresponding number of samples n_d for direct learning.

The significant literature on direct learning indicates that various

methods may be used to decide the n_d . One approach is to fix an arbitrary error limit and incrementally increase the number of experimental training points generated till the testing error reduces below this limit. This approach is fraught with an arbitrary subjectivity of how much error is acceptable. The second approach is to a-priori designate the maximum number of experimental points that can be generated and take the corresponding testing error at this maximum limit as the lowest possible error. This method is subject to unnecessarily increasing the amount of experimental data needed, e.g., the error may not reduce significantly with the addition of experimental data beyond a certain threshold. The third approach involves setting an experimental budget N_B and incrementally increasing the number of experimental training points till the percentage difference in testing error across consecutive increments goes below a user-defined threshold ζ , or there is overfitting, or the number of training points exceeds N_B , whichever occurs first. The idea is that if the error difference from one increment to the next goes below ζ then there is no point in adding further experimental points since it will not realize a significant enough error reduction. This approach enables a more objective balance between the amount of experimental data needed and the error while preventing overfitting. In this work the third approach was used and the ζ was fixed as 2 %. The process-specific N_B was based on the experimental capabilities in the author's laboratory, as would be the case in industrial practice, and is mentioned with the description of the experimental methods.

For Smart-ML, a separate ε -SVR was trained on the data from the source process model. An incrementally increasing amount of target data was used to iteratively identify the smallest amount of experimental data needed for transfer learning (n_t) while satisfying the constraint that the error of the final ε -SVR δ_t on a withheld experimental dataset should be lesser than or equal to the above computed δ_d . This constraint ensured that the drive towards reducing C_{DEV} did not compromise the prediction accuracy. Note that this paper does not focus on the method for sampling the experiments since its emphasis is on showing the feasibility of Smart-ML and the corresponding cost savings that are possible. The test dataset was of the same size as that corresponding to n_d to prevent a heavily lopsided train:test ratio and thus fairly compare direct learning and Smart-ML. It was obtained randomly from the withheld target dataset. The aforementioned testing for the final ML model was performed thirty times over the entire dataset to obtain the mean and standard deviation

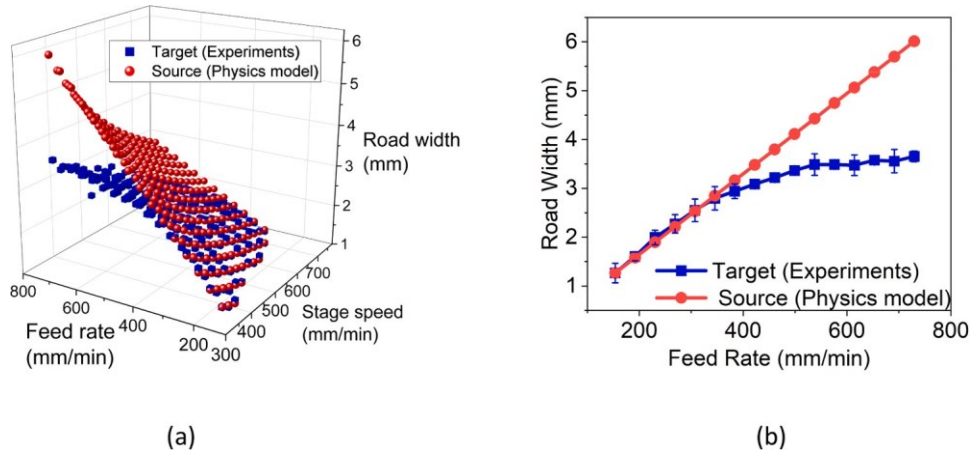


Fig. 5. Comparison of source and target (a) 3D plot (b) 2D plot at a constant stage speed of 350 mm/min. All shown for a nozzle-to-stage distance of 0.7 mm.

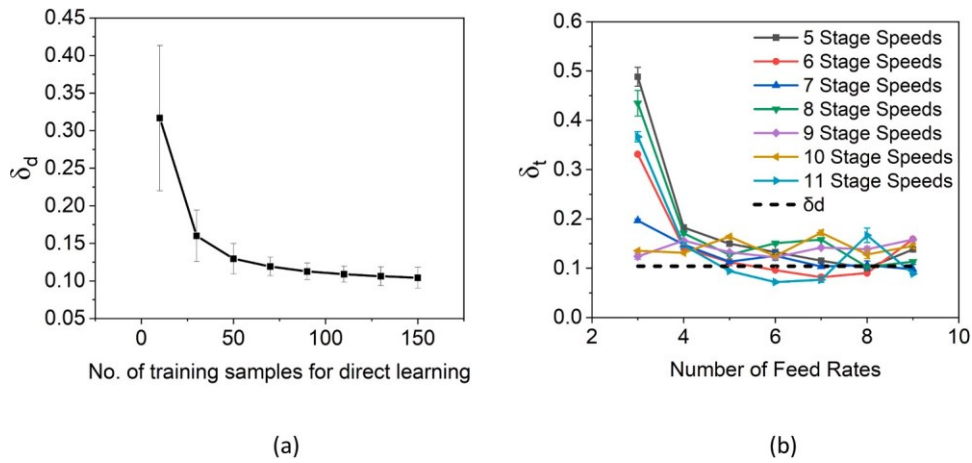


Fig. 6. Change in error on testing dataset with (a) number of total training points in direct learning (b) combinations of number of stage speeds and feed rates in Smart-ML.

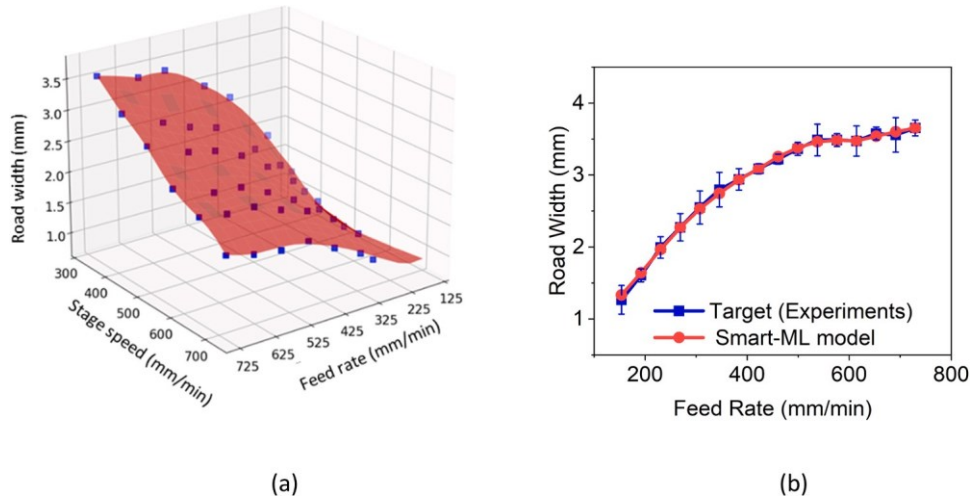


Fig. 7. Comparison of final ML model after Smart-ML to experimental target (a) 3D plot. The red surface shows the final ML model. The blue squares represent the target data (b) 2D plot at stage speed of 350 mm/min. All shown for a nozzle-to-stage distance of 0.7 mm.

of δ_t . Note that overfitting was not observed in this work for either direct learning or Smart-ML since the testing error did not exceed the training error.

2.2. Source modeling and experimental methods

The capabilities of Smart-ML were explored for three manufacturing problems. This section describes these processes and the source models

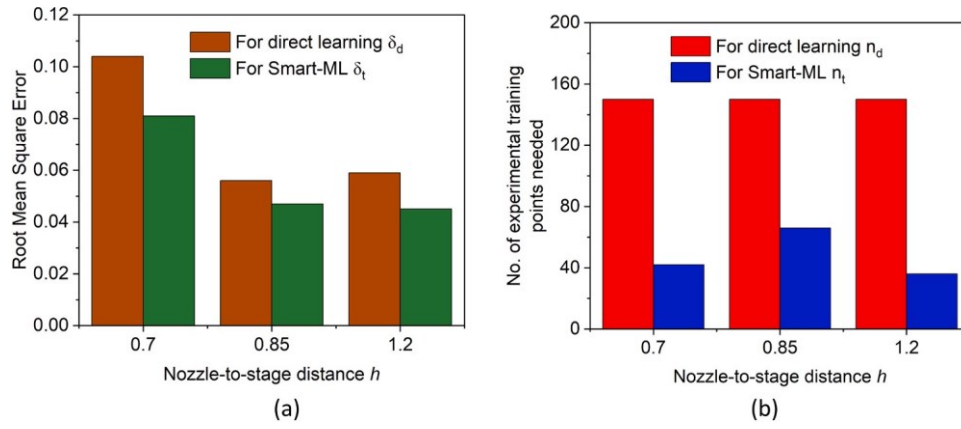


Fig. 8. Comparison of direct learning and Smart-ML in terms of (a) error on testing dataset (b) number of experimental samples needed for training.

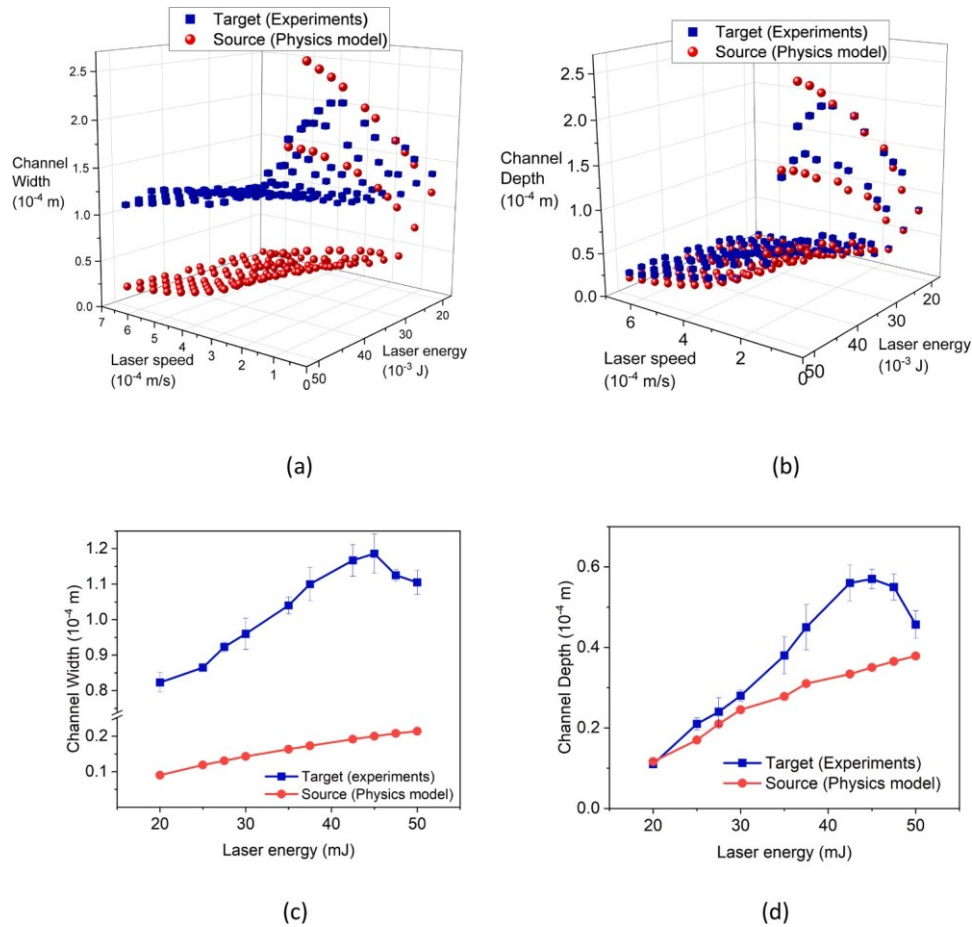


Fig. 9. Comparison of source and target for prediction of microchannel (a) width (b) depth. Representative 2D plots of (c) Microchannel width at laser speed of 650 μm/s (d) Microchannel width at laser speed of 350 μm/s. All shown for a laser pulse frequency of 5 Hz.

used.

2.2.1. Fused Filament Fabrication (FFF)

The first problem, for an additive process, involves prediction of the road width printed in Fused Filament Fabrication (FFF) as a function of the stage speed and the filament feed rate (Fig. 4a). This problem has been tackled over the last two decades with incrementally increasing accuracy, including most recently by using computational fluid dynamics models (Serdeczny et al., 2018). These efforts have revealed that the underlying physics involves non-Newtonian flow, compressibility,

nozzle-extrudate interaction, wetting and non-isothermal cooling.

The source process model used here is $W = FA/Sh$ where W is the road width, F is the filament feed rate, S is the stage speed, A is the filament's cross-sectional area, and h is the nozzle-to-stage distance. This source assumes that the above physics and the corresponding constitutive laws are unknown, makes the incorrect but simplifying assumption that the nozzle-to-stage distance equals the height of the road, and includes the mass conservation law.

The experimental target data was generated by printing roads with a 1 mm diameter nozzle, nozzle temperature of 230 °C, and bed

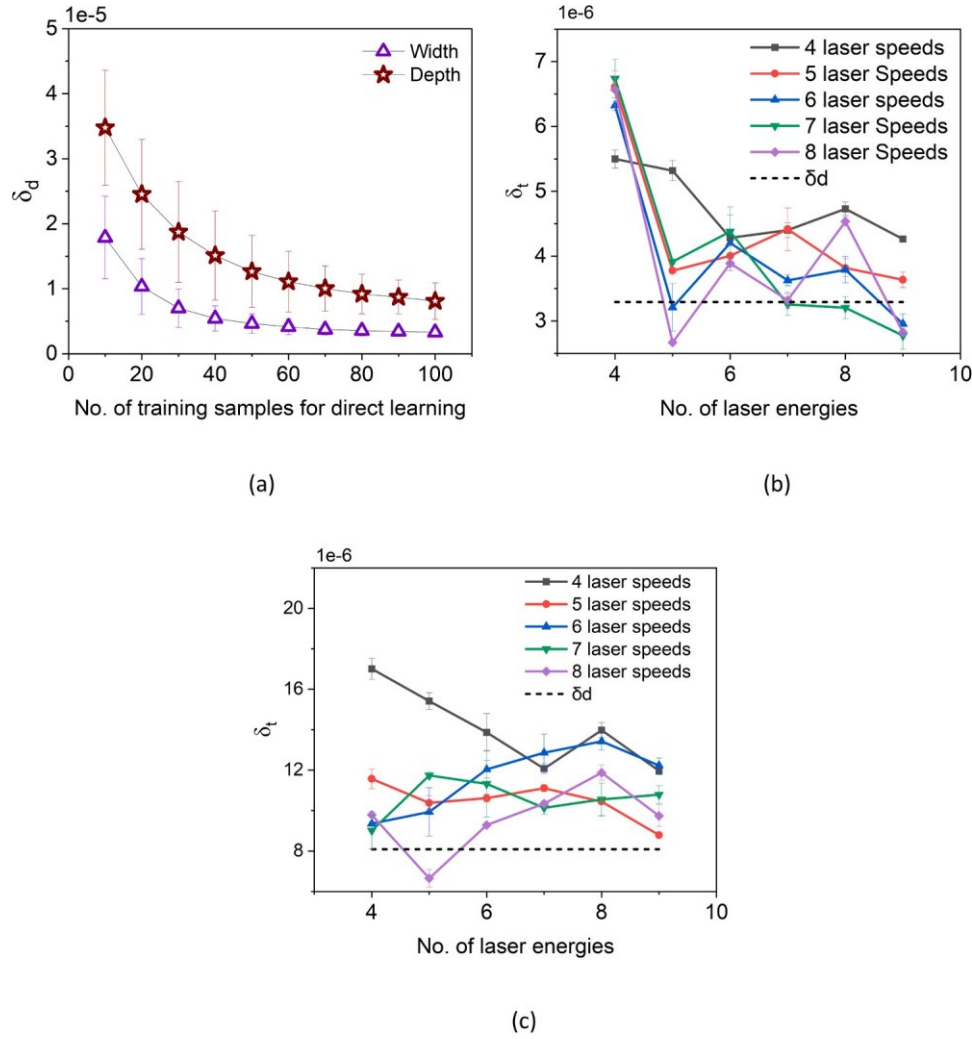


Fig. 10. (a) Direct learning testing error with number of total training points. Smart-ML testing error with different combinations of number of experimental laser speeds and energies for (b) Channel width and (c) Channel depth. All for 5 Hz laser frequency.

temperature of 60 °C. A PLA filament of 0.75 mm diameter (from 3DExTech) without any particulate additives was used. For each value of h explored here the experimental data was generated across 256 distinct combinations of S and F . Specifically, sixteen equidistant S (between 350 and 725 mm/min) and F (between 153 and 729 mm/min) were used. Three different h values of 0.7, 0.85, and 1.2 mm were explored. The W was measured using vernier calipers and unstable printing regimes were excluded from the dataset. The experimental budget N_B for each h was fixed at a maximum of 150 training points.

Since this problem has an established and accurate physics-based model that is applicable across the wide parameter range used in our experiments, this process testbed allows decisive quantification of the reduction in C_{DEV} achieved by Smart-ML. The C_{DEV} is computed as the time in human-years between the first public report of the source model used for Smart-ML and the accurate physics-based model. The savings in C_{EXP} are calculated as the percentage change between n_d and n_t , as in the multifidelity learning literature. The savings in C_{COMP} are based on the difference between the total computational effort (in CPU-hours) for the accurate physics-based model, as reported in the literature, and the total CPU-hours for the source model used for Smart-ML. Since these savings only make sense if enough predictive accuracy is achieved using Smart-ML they are reported in detail along with the testing error in the results section.

2.2.2. Magnetically-Assisted Laser Induced Plasma Micromachining (M-LIPMM)

The second problem, for a subtractive process, was to predict the microchannel width and depth as a function of the laser speed and energy in the newer Magnetically Assisted Laser Induced Plasma Micromachining (M-LIPMM) process. LIPMM, illustrated in Fig. 4b, creates a plasma in a dielectric liquid and then uses it for machining (Pallav et al., 2015). M-LIPMM builds on the low thermal damage and optically-agnostic material capabilities of LIPMM by adding a magnetic field that manipulates the plasma and the machined feature's dimensions beyond the optical and thermal limitations of LIPMM (Malhotra et al., 2013). Several physical phenomena have been hypothesized as being relevant to material removal in M-LIPMM. These include dielectric breakdown, magnetohydrodynamics of evolution in plasma size, pressure, and temperature, plasma-workpiece interactions via thermal ablation and athermal material removal, heat transfer to the fluid and bubble formation, and pressure-induced expulsion of the ablated material between consecutive laser pulses. (Zhang et al., 2021) discussed the competing effects of magnetic fields and confinement-induced absorption on the plasma size, temperature, and pressure. But the lack of plasma-workpiece interactions, the simplified laser-plasma interaction, and the qualitative form of the model, meant that the feature size could not be predicted. (Saxena et al., 2014) and (Xie et al., 2020) modeled plasma generation due to dielectric breakdown but without a magnetic field or plasma-workpiece interactions.

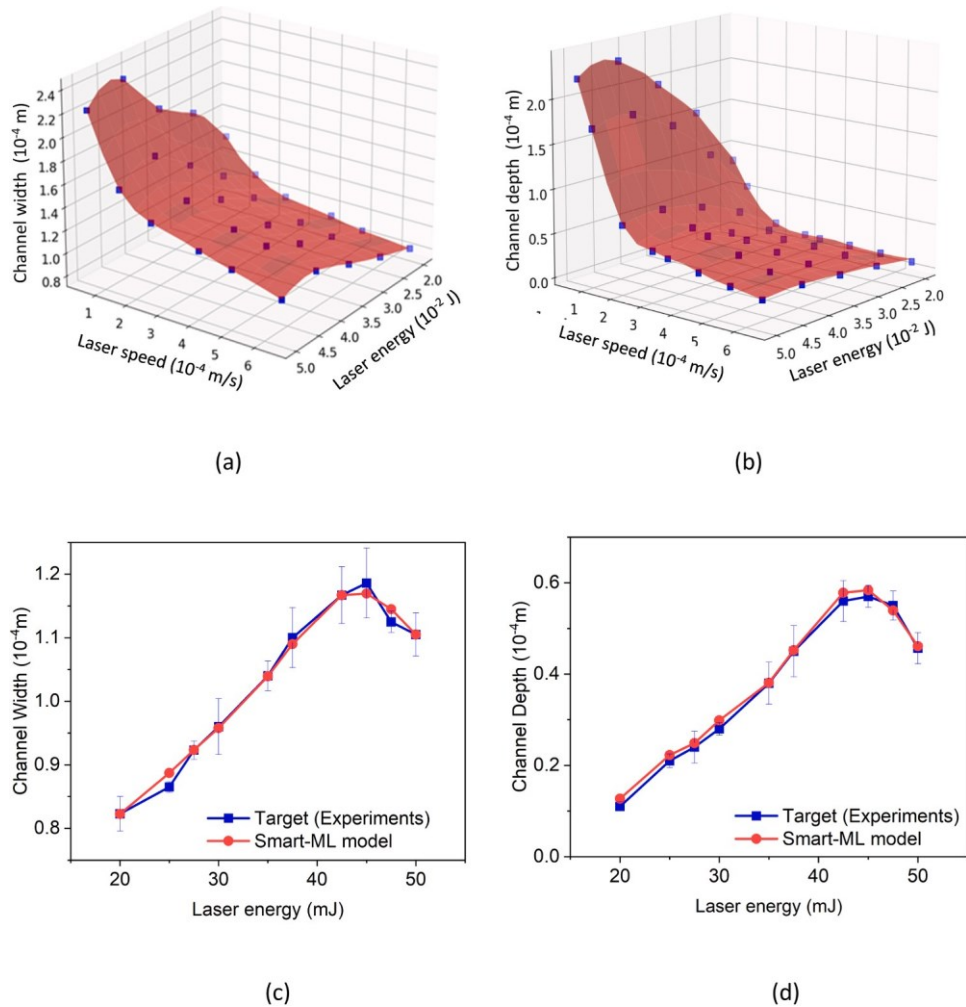


Fig. 11. Comparison of final Smart-ML model to experimental target (a-b) 3D plots for channel width and depth. The red surface shows the final ML model. Blue squares represent target data. Corresponding 2D plots for microchannel (c) width at laser speed of 650 $\mu\text{m/s}$ (d) width at laser speed of 350 $\mu\text{m/s}$. All shown for a 5 Hz pulse frequency.

(Wang et al., 2018) modeled plasma evolution in the absence of magnetic effects or plasma-workpiece interactions resulting in high prediction errors of up to 50 %. Thus, purely physics-based modeling efforts over the last decade cannot quantitatively model the machined feature size in M-LIPMM. Thus, this problem is a true test of the capabilities of Smart-ML since significant components of the process physics are currently unknown and cannot be explicitly modeled.

The source process model in this paper ignores significant components of the above hypothesized and modeled physics. It models material removal based on a 2D Finite Element Analysis (FEA) of thermal conduction in the workpiece within the COMSOL platform (Fig. 4d). This accounts for the conservation of energy, mass and momentum but with the following simplifications that omit the derivation of multiple multiphysical couplings and constitutive models of material properties. The heat source that represented the plasma heat flux into the workpiece had a uniform spatial distribution over the region corresponding to the diameter of the laser's focal spot (see Fig. 4d) and was non-zero only during the pulse duration of the laser (6 ns here). Complete absorption of the laser pulse energy by the workpiece was assumed, thus eliminating any thermal effects of the liquid and the optical-thermal-hydrodynamic coupling inherent to the plasma formation and its interaction with the workpiece. The laser was assumed to be stationary to increase the computational speed of the FEA. Thus, thermal conduction into the workpiece during only one laser pulse was modeled. The channel depth and width created in this one laser pulse was

manually extracted based on the spatial location of the elemental integration points at which the workpiece's melting point was exceeded. The laser's residence time at a spot was calculated analytically based on its speed and optics-based spot size. The corresponding number of pulses at a given point was thus obtained based on the pulse frequency, as is common in direct laser micromachining with low pulse frequencies (Ling, 2011). The machined channel's depth and width was predicted as the product of the dimensions created by one pulse with the above calculated number of pulses at a given material point. The underlying assumption behind this method of calculating channel depth and width is that material removal is purely additive over multiple pulses. Constant thermal properties of the aluminum workpiece were assumed. The density was fixed at 2.7 g/cm^3 , the thermal conductivity at 237 W/m-K , the specific heat capacity at 903 J/kg-K , and the melting point at 934 K. This state-independent assumption is especially incorrect at the elevated workpiece temperatures expected in M-LIPMM. Further, these properties were not measured but were based on commonly reported values in the literature. The boundary conditions consisted of insulation at all boundaries except where the laser flux was applied (Fig. 4d). The initial condition was room temperature for the workpiece. The workpiece width and depth were 200 μm each, i.e., larger than the laser spot of 80 μm used in experiments. Square mesh elements were used with a uniform side length of 5 μm . No tests for mesh convergence or element type were performed. Explicit time-marching simulations were performed with the time steps decided automatically by the COMSOL

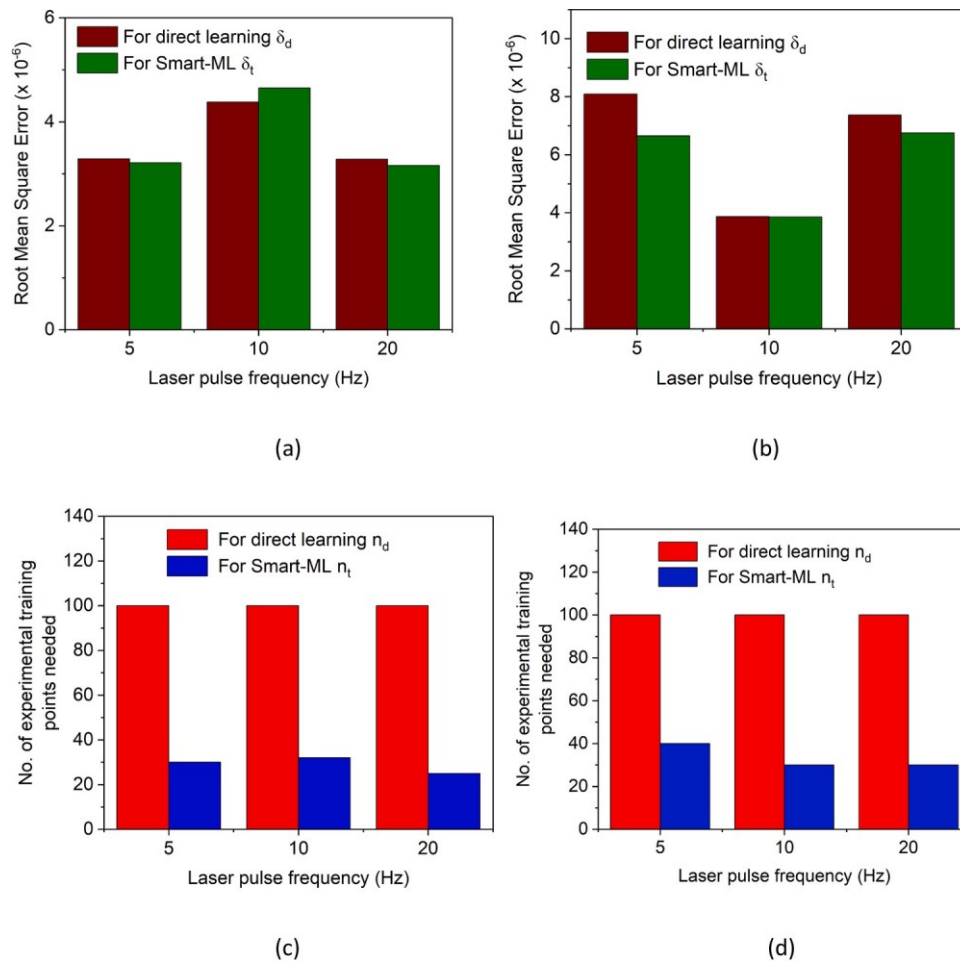


Fig. 12. Comparison of direct learning and Smart-ML in terms of error on testing dataset for prediction of (a) channel width (b) channel depth; and number of experimental samples needed for training for (c) channel width (d) channel depth.

software without user manipulation. Fig. 4d shows a representative example of the temperature evolution predicted by this FEA.

Since this problem does not have an established and accurate physics-based model, but the above simple source model has always been available, the savings in C_{DEV} is a lower bound estimate based on the time elapsed since the M-LIPMM process was first publicly reported in 2013. The savings in C_{EXP} are calculated as in the FFF problem. The savings in C_{COMP} cannot be directly computed since the accurate physics-based model is not available. Instead, we compare the CPU-hours needed to generate the source data in Smart-ML to that needed for computation of whatever sub-components of the process physics can be accurately modeled today, namely the spatiotemporal evolution of the size of the plasma under a magnetic field without modeling plasma-material interactions (Kuzenov and Ryzhkov, 2018). These savings are quantified in detail in the results section in light of the achieved prediction accuracy.

The M-LIPMM experiments were performed on a Quantum Light Instruments pulsed laser (3 ns pulse duration, 50 mJ pulse energy, 526 nm wavelength). The 4 mm diameter laser beam was optically reduced to a non-diffraction-limited 80 μ m diameter at the focal point where the plasma was produced. The workpiece material was 6061 Aluminum from McMaster Carr. The dielectric liquid was deionized water. The workpiece was kept in a glass beaker with a 2 mm thick layer of water above it. The glass beaker was, in turn, kept on a motorized XYZ stage that had a resolution of 50 nm. The microchannels were machined by moving the laser at a specified speed along the desired direction in a single pass, i.e., without translating the laser spot in the vertical

direction. The dimensions of the machined channels were measured using a Keyence optical profilometer. This experimental data was generated for combinations of thirteen equidistant laser speeds (5×10^{-5} to 6.5×10^{-4} m/s) and 10 equidistant laser energies (from 20 to 50 mJ), over each of the three laser pulse frequencies (5, 10 and 20 Hz). Thus, the total number of experimental training points for each pulse frequency was 130. The experimental budget N_B for each pulse frequency was fixed at 100 so that enough withheld testing points were also available.

2.2.3. In-situ rolling in fused filament fabrication

The third problem, for a hybrid additive-deformation process, was in-situ rolling in FFF. In this heated and weighted rollers attached to the FFF extruder heat and compress the printed road to eliminate voids between roads (Fig. 4c). This kind of hybrid additive-deformative process for polymer extrusion additive manufacturing was first proposed by (Duty et al., 2017). It markedly increases the strength and isotropy of FFF parts. The modeling problem is to predict the part height induced by in-situ rolling as a function of the weight on the roller and the nozzle-to-stage distance. The only physics-based solution till date for this problem is an analytical model of the contact width between two cylinders (representing the roads) that are subjected to an applied load (Qasimeh et al., 2022). This is contingent on the assumption that the as-deposited road has a circular cross-section with a diameter equal to the nozzle diameter, thus ignoring the nozzle-extrudate interactions that cause as-printed roads to have non-circular cross sections and greater width than the nozzle diameter (Serdeczny et al., 2018). This model is

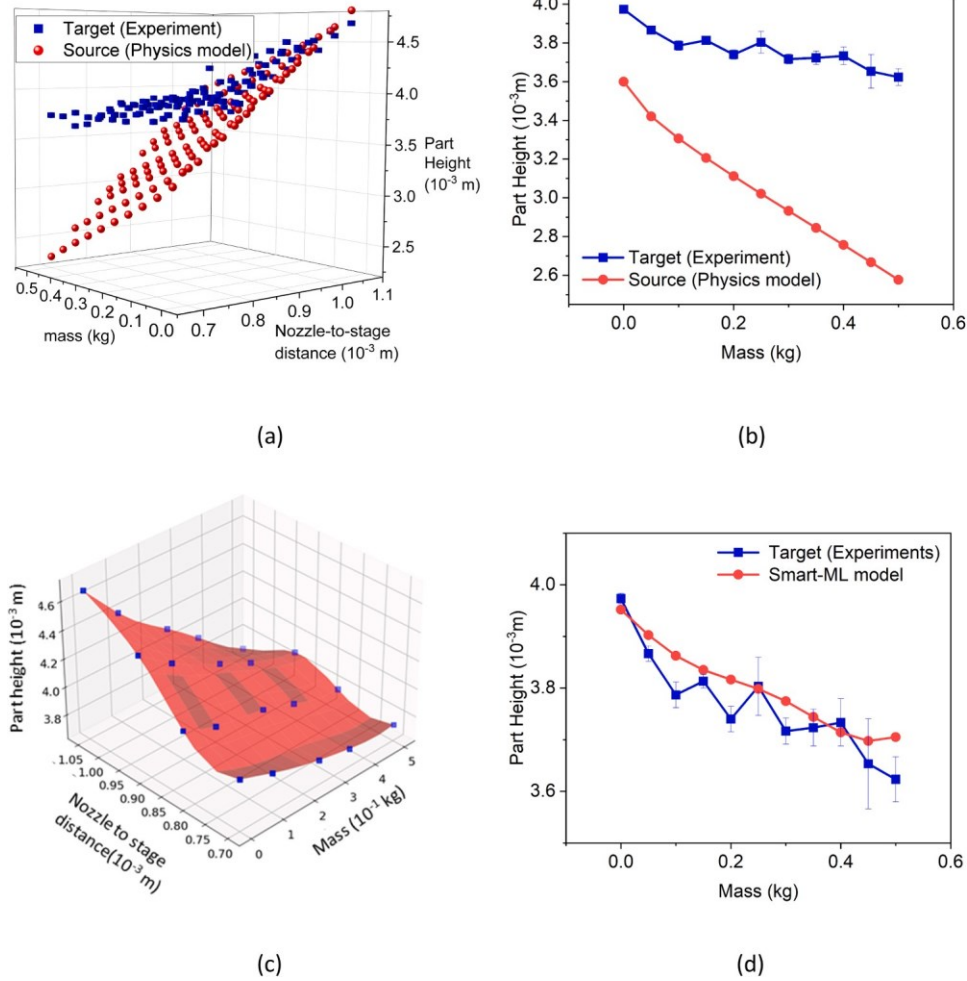


Fig. 13. Comparison of source and target (a) 3D plot (b) 2D plot at a constant nozzle-to-stage distance of 0.75 mm. Comparison of final Smart-ML model to experimental target (c) 3D plots. The red surface shows the final ML model. The blue squares represent the target data. (d) 2D plot at a constant nozzle-to-stage distance of 0.75 mm.

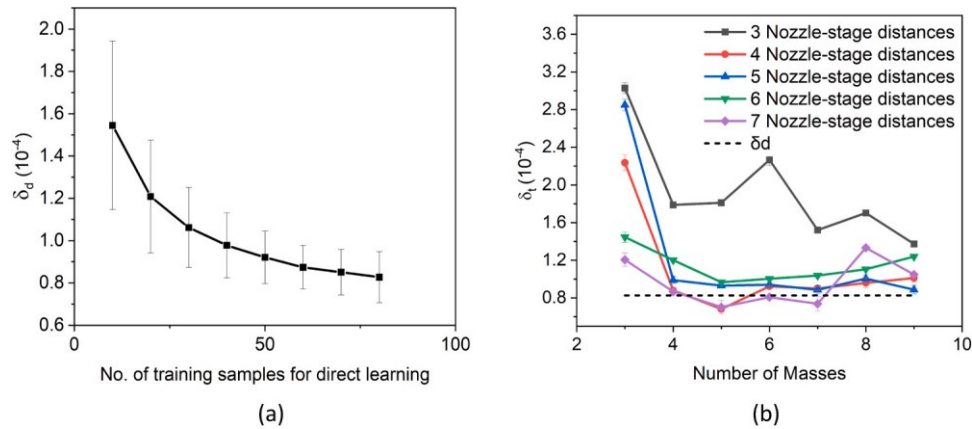


Fig. 14. Change in testing error with (a) number of training points used in direct learning (b) combinations of number of masses and nozzle-to-stage distances used in Smart-ML.

consequently unable to consider the convolution between the effect of layer height and roller mass. This problem constitutes a case where two different physical effects in a hybrid process are known, i.e., compression by the rollers and nozzle-extrudate interaction, but a model that couples these phenomenon is as yet unavailable. This paper uses the above analytical model from (Qasaimeh et al., 2022) as the source

process model. Since this problem also does not have an established and accurate physics-based model the savings in C_{DEV} , C_{EXP} and C_{COMP} are computed as in the case of the M-LIPMM problem and are quantified in detail in the results section.

Experiments were performed on a commercial FFF printer after the approach described by (Qasaimeh et al., 2022). PLA polymer was used

with a nozzle temperature of 230 °C, bed temperature of 60 °C, roller temperature of 110 °C, roller diameter of 6.35 mm, and nozzle diameter of 1.8 mm. The change in part height over four consecutive layers was measured using vernier calipers. The experiments were performed for combinations of ten equidistant roller weights (0–500 g) and ten equidistant nozzle-to-stage distances (0.70–1.5 mm). Thus, the total number of experimental points generated was 100. The experimental budget N_B was fixed at a maximum of 80 training points so that sufficient withheld testing points were also available.

3. Results

3.1. The FFF problem

Fig. 5a shows a representative comparison between the experimentally measured (target) road width and the predictions of the source process model for the FFF problem. The source can only make good predictions at lower feed rates and stage speeds. The extent of the qualitative difference is seen in the representative 2D plot in Fig. 5b. A quantitative recalibration of the source model or scaling of the source data cannot capture this experimentally observed parametric effect, and a functional correction is needed.

Fig. 6a shows an example of the change in δ_d with the number of training samples used for direct learning, yielding a n_d of 150. The corresponding evolution of the transfer learning error δ_t for distinct combinations of the number of stage speeds and filament feed rates used from the target dataset is reported in Fig. 6b. The minimum number of experimental training points (n_t) for which $\delta_t \leq \delta_d$ is 42, nearly 72 % lesser than n_d . The corresponding reduction in the testing error with Smart-ML relative to direct learning was 22 %. The ability of transfer learning to capture the experimental observations both qualitatively and quantitatively is illustrated via the examples in Fig. 7. Fig. 8 shows the extent of these advantages for the other h values used here.

A reduction in C_{EXP} and error is also seen in the literature on multi-fidelity learning in manufacturing, but not for the kind of functional differences between the source and the target used here. This shows that the current advantages of multi-fidelity learning can be retained if the source uses the conservation equations, even with functionally incomplete or quantitatively inaccurate constitutive laws, to capture at least some of the experimental trends albeit in a qualitatively incorrect manner. For example, the source captures the fact that the road width reduces with increasing stage speed but in a functionally incorrect way. This allowance for a qualitative discrepancy between the source and target is critical for achieving low C_{DEV} for new processes. Existing functionally accurate and physically complete process models can be used as the source for multi-fidelity learning. But it has taken significant incremental human effort over long periods of time for these models to be developed to this level of functional accuracy. The source model used here was first proposed in the year 2000 at the dawn of research into road formation in FFF (Bellini et al., 2004). Analytical physics-based models have taken till 2019 to predict the experimentally observed nonlinearity over an appreciably wide range of filament and stage speeds, specifically the model by (Agassant et al., 2019) based on a 3D approximation of 2D Stokes flow with simple shear and isothermal wetting. Computational fluid dynamics simulations, which started in 2002 with the dissertation of (Bellini, 2002), have taken till the 2018 work of (Comminal et al., 2018) to predict road width across a substantial range of filament and stage speeds. Thus, using Smart-ML in 2000, which is when the employed source model was available, could have saved at least 18 years of C_{DEV} .

Finally, the authors note that using computational fluid dynamics models to generate just one training sample for FFF needs orders of magnitude more CPU-hours than the 3×10^{-6} CPU-hours needed here to generate the entire source dataset for Smart-ML. This is because a simplified analytical process model with no spatiotemporal discretization is used as the source. While the widespread advent of significant

computational capacity might make this advantage moot, it is worth noting its existence.

3.2. The M-LIPMM problem

The experimentally measured channel width and depth in Fig. 9a-b exhibit an inflection point at high laser energy and low laser speed that is not seen in the source. This inflection is even clearer when plotting the channel width and depth as a function of the laser energy for a fixed laser speed, as shown in Fig. 9c-d.

Fig. 10a reports the evolution of δ_d with the number of experimental training samples. The corresponding n_d was 100. Fig. 10b-c show that Smart-ML can achieve $\delta_t \leq \delta_d$ with a 60–70 % reduction in the number of experimental training samples, and thus in C_{DEV} . The ability of the final ML model to capture the experimentally observed inflection point is clear in Fig. 11. Compared to the FFF problem where the target data is monotonic, the ability to capture the inflection point in M-LIPMM is a more extreme example of the degree to which Smart-ML can compensate for qualitative discrepancies between the source and the target. Fig. 12 shows that reductions in C_{EXP} are possible across different laser pulse frequencies without compromising the error significantly.

There is currently no physics-based model that can predict the experimentally observed non-monotonic response of the channel dimensions in M-LIPMM to the laser energy and speed. Given that the M-LIPMM process was first reported a decade ago in (Malhotra et al., 2013), and the simple source model used in this paper has been available at least since then if not earlier, the use of Smart-ML could have saved at least 10 human-years of C_{DEV} . Further, just one simulation of the plasma-magnet interaction even without modeling plasma-material interaction needs at least 10 times more CPU-hours (Kuzenov and Ryzhkov, 2018) than the total C_{COMP} of 2 CPU-hours to generate the source data in Smart-ML here. Thus, the reduction in C_{COMP} is also significant. Overall, these results show that even large functional differences between the source and the target, represented by a monotonic source and a non-monotonic target here, can be accommodated by Smart-ML while retaining high reductions in C_{DEV} , C_{EXP} , and C_{COMP} and high prediction accuracy.

3.3. The in-situ rolling problem

Fig. 13a-b show the difference between the linear source model and the nonlinear target data which leads to significant differences in the predicted part height at a smaller layer height. Smart-ML compensates for this discrepancy by accurately capturing the experimentally observed convolution of the layer height and roller mass effects on the part height (Fig. 13c-d).

The evolution of the error for direct learning and Smart-ML for different number of training samples, as shown in Fig. 14a-b, yields an n_d of 80 and a n_t of 20 (i.e., 4 layer heights and 5 roller masses). This corresponds to a 17 % reduction in the error from a δ_d of 8.2×10^{-5} to a δ_t of 6.8×10^{-5} and an 75 % reduction in C_{EXP} . As described earlier, this in-situ rolling process was introduced in 2017 and the source process model used here was developed in 2022. Assuming a similar timeframe for advancing this process model by coupling past computational fluid dynamics models of road formation (Serdeczny et al., 2018) with thermomechanical finite element models of deformation of roads with a non-circular cross section by two hot rollers, it is estimated that Smart-ML can save at least 5 human-years of C_{DEV} . Further, the C_{COMP} of 10^{-6} CPU-hours needed to generate all the source samples here is expected to be orders of magnitude lower than the above advanced model. Thus, Smart-ML can correct process models which may be only valid in certain regions of the parameter space due to the absence of a coupling between critical physics.

4. Conclusions

This paper develops and demonstrates a Smart-ML approach for machine learning of parametric relationships in manufacturing processes, especially new ones, for which a qualitative understanding of the process physics, constitutive relationships, and multiphysical and multiscale couplings is either unavailable or incomplete. The key advance is a reduction not just of the experimental and computational costs but also of the often-ignored and significant physics development cost of generating data for training the ML models. The three process examples for which Smart-ML is explored show that the reduction in the physics development cost can be on the order of multiple human-years, the reduction in experimental costs can be as high as 60 %, and the reduction in computational costs could be on orders of magnitude, without an increase in prediction error. This demonstration across three processes with different physical principles and different levels of physics-based understanding builds confidence that Smart-ML can be used for rapid, accurate, and experimentally inexpensive training of ML models of parametric relationships in manufacturing processes.

The reduction in C_{DEV} is achieved by relaxing the constraint that the source and the target in multi-fidelity transfer learning should be functionally similar. It is possible to retain low prediction error despite this greatly simplifying restriction on the fidelity of the source models. This is because the mandatory inclusion of the conservation equations in the source model provides a physical regularization, albeit incomplete, that leads the ML model partially towards the ground truth. Since the experimental data inherently adheres to the conservation laws the transfer learning-based update of the ML model compensates for the other simplifications made in the source model. The key insight is that this regularization need not be qualitatively complete, unlike the assumption in the state-of-the-art, which enables significant reductions in the physics development cost by eliminating the need to iteratively calibrate and validate the physics-based source model.

The reduction in C_{EXP} is possible because the source data satisfies some of the hunger of the ML model for data. Specifically, the source data allows an initial estimation of the weights of the base ML model (ϵ -SVR here) so that the amount of experimental data needed for the finalizing the ML model's weights to match predictions to the ground truth is lesser than that needed for direct learning. Note that the approach for choosing n_d for direct learning balances the number of experimental training points needed and the error. It is also worth noting that the experimental budget for direct learning N_B may be reduced arbitrarily, but at the cost of significant error and without any significant advantage over Smart-ML. For example, fixing the N_B and thus n_d at 50 in the case of FFF increases the RMSE by nearly 25 % as compared to a n_d of 150 points (Fig. 6a). Further, the reduction in error from 30 to 50 experimental training points for direct learning is a significant 20 % (Fig. 6a). This indicates that more experimental training points are certainly needed for direct learning. Finally, Smart-ML still achieves a lower error of 0.123 with 27 experimental training points (3 feed rates and 9 stage speeds in Fig. 6b) which constitutes a 46 % reduction in C_{EXP} even with $n_d = 50$. Continuing with the same example, one might say that the N_B and thus the n_d could be reduced even further to 30, but the error will be sacrificed too significantly to claim that the direct learning-based ML model has any reasonable accuracy even compared to Smart-ML with 3 feed rates and 9 stage speeds (Fig. 6b).

The reduction in C_{COMP} is realized because Smart-ML enables the use of simplified physics-based models (i.e., source) which need not be quantitatively accurate or qualitatively complete. This enables significant reductions in the amount of multiphysical and multiscale couplings needed in the physics-based model, e.g., complete omission of the plasma-magnet coupling in the M-LIPMM, and significantly reduces the amount of spatiotemporal discretization needed, e.g., a simple analytical model without finite element discretization or time-marching solvers can be used in the FFF problem. Meanwhile, the increase in error due to the reductions in quantitative and functional accuracy of the physics-

based model is compensated for by the transfer learning component of Smart-ML.

Training ML models of parametric relationships on model-generated rather than experimentally-generated data has become increasingly common to reduce the experimental cost. But this approach results in ML being restricted to processes for which physics-based process models are mechanistically complete. The above advantages of Smart-ML can change this paradigm by accelerating ML of processes even when there is significantly incomplete or limited knowledge on their underlying physics. It is envisioned that Smart-ML will transform the function of ML from the calibration and replication of existing physics-based models in a computationally efficient form to rapid model discovery for difficult to model or novel processes. At the same time, Smart-ML is not meant to replace human intuition but to augment it. The authors believe that the critical role of the process engineer in Smart-ML of choosing the source model will create a new form of human-machine interaction in the realm of manufacturing process modeling.

But there are also challenges in the current form of Smart-ML. The first is the lack of appropriate methods for sampling the target so that the experimental cost can be reduced in practice by identifying the target dataset in an a-priori and intelligent manner. The second is an effort to extend Smart-ML to modeling evolution of the spatiotemporal material state in manufacturing processes where using advanced base ML models will be necessary. These areas of future work are being pursued by the authors. Finally, the process testbeds used in this work are limited by the resources in the author's laboratory. The authors anticipate greater testing of Smart-ML across multiple additional processes in the larger manufacturing processes community and welcome such collaborations.

CRedit authorship contribution statement

Jeremy Cleeman: Methodology, Software, Validation, Data curation, Writing- Original Draft, Visualization **Kian Agrawala:** Methodology, Software, Validation, Data curation, Writing- Original Draft **Evan Nastarowicz:** Investigation, Data curation **Rajiv Malhotra:** Conceptualization, Resources, Writing- Review and Editing, Project Administration, Funding Acquisition.

Declaration of Competing Interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Data Availability

Data will be made available on request.

Acknowledgements

This work was supported by the National Science Foundation, USA, grant CMMI # 2001081.

References

- Agassant, J.-F., Pigeonneau, F., Sardo, L., Vincent, M., 2019. Flow analysis of the polymer spreading during extrusion additive manufacturing. *Addit. Manuf.* 29, 100794.
- Alam, M.F., Shtein, M., Barton, K., Hoelzle, D.J., 2020. Autonomous manufacturing using machine learning: A computational case study with a limited manufacturing budget. *International Manufacturing Science and Engineering Conference* 84263, V002T007A009.
- Arinez, J.F., Chang, Q., Gao, R.X., Xu, C., Zhang, J., 2020. Artificial Intelligence in advanced manufacturing: current status and future outlook. *J. Manuf. Sci. Eng.* 142.
- Barnes, A.J., 2007. Superplastic forming 40 years and still growing. *J. Mater. Eng. Perform.* 16, 440–454.
- Bellini, A., 2002. Fused deposition of ceramics: a comprehensive experimental, analytical and computational study of material behavior, fabrication process and equipment design. PhD dissertation, Drexel University, Philadelphia, PA.

- Bellini, A., Gucci, S., Bertoldi, M., 2004. Liquefier dynamics in fused deposition. *J. Manuf. Sci. Eng.* 126, 237–246.
- Cleeman, J., Bogut, A., Mangrolia, B., Ripberger, A., Kate, K., Zou, Q., Malhotra, R., 2022. Scalable, flexible and resilient parallelization of fused filament fabrication: Breaking endemic tradeoffs in material extrusion additive manufacturing. *Addit. Manuf.* 56, 102926.
- Comminal, R., Serdeczny, M.P., Pedersen, D.B., Spangenberg, J., 2018. Numerical modeling of the strand deposition flow in extrusion-based additive manufacturing. *Addit. Manuf.* 20, 68–76.
- Daehn, G.S., Taub, A., 2018. Metamorphic manufacturing: the third wave in digital manufacturing. *Manuf. Lett.* 15, 86–88.
- Devaraj, H., Malhotra, R., 2019. Scalable forming and flash light sintering of polymer-supported interconnects for surface-conformal electronics. *J. Manuf. Sci. Eng.* 141.
- Devaraj, H., Tian, Q., Guo, W., Malhotra, R., 2021. Multiscale modeling of sintering-driven conductivity in large nanowire ensembles. *ACS Appl. Mater. Interfaces* 13, 56645–56654.
- Drucker, H., Burges, C.J., Kaufman, L., Smola, A., Vapnik, V., 1996. Support vector regression machines. *Adv. Neural Inf. Process. Syst.* 9.
- Duflou, J.R., Habraken, A.-M., Cao, J., Malhotra, R., Bambach, M., Adams, D., Vanhove, H., Mohammadi, A., Jeswiet, J., 2018. Single point incremental forming: state-of-the-art and prospects. *Int. J. Mater. Form.* 11, 743–773.
- Duty, C.E., Kunc, V., Compton, B., Post, B., Erdman, D., Smith, R., Lind, R., Lloyd, P., Love, L., 2017. Structure and mechanical behavior of Big Area Additive Manufacturing (BAAM) materials. *Rapid Prototyp. J.*
- Fan, Z., Ho, J.C., Takahashi, T., Yerushalmi, R., Takei, K., Ford, A.C., Chueh, Y.L., Javey, A., 2009. Toward the development of printable nanowire electronics and sensors. *Adv. Mater.* 21, 3730–3743.
- Hounshell, D.A., 1984. *From the American System to Mass Production, 1800–1932: The Development of Manufacturing Technology in the United States.* Johns Hopkins University Press Baltimore, Maryland.
- Huang, S.H., Liu, P., Mokasdar, A., Hou, L., 2013. Additive manufacturing and its societal impact: a literature review. *Int. J. Adv. Manuf. Technol.* 67, 1191–1203.
- Huang, X., Xie, T., Wang, Z., Chen, L., Zhou, Q., Hu, Z., 2021. A transfer learning-based multi-fidelity point-cloud neural network approach for melt pool modeling in additive manufacturing. *ASCE-ASME J. Risk Uncert. Engrg. Sys. Part B Mech. Engrg.* 8.
- Jain, R.K., Smith, K.M., Culligan, P.J., Taylor, J.E., 2014. Forecasting energy consumption of multi-family residential buildings using support vector regression: investigating the impact of temporal and spatial monitoring granularity on performance accuracy. *Appl. Energy* 123, 168–178.
- Kapusuzoglu, B., Mahadevan, S., 2020. Physics-informed and hybrid machine learning in additive manufacturing: application to fused filament fabrication. *JOM* 72, 4695–4705.
- Kuzenov, V.V., Ryzhkov, S.V., 2018. Numerical modeling of laser target compression in an external magnetic field. *Math. Models Comput. Simul.* 10, 255–264.
- Ling, T.D., 2011. *PhD thesis: Mechanics and Control of Laser Surface Texturing and its Applications in Energy Efficiency and Production, Mechanical Engineering.* Northwestern University, Evanston, Illinois.
- Liu, Z., Guo, B., Wu, F., Han, T., Zhang, L., 2022. An improved burr size prediction method based on the 1D-ResNet model and transfer learning. *J. Manuf. Process.* 84, 183–197.
- Lockner, Y., Hopmann, C., 2021. Induced network-based transfer learning in injection molding for process modelling and optimization with artificial neural networks. *Int. J. Adv. Manuf. Technol.* 112, 3501–3513.
- Lockner, Y., Hopmann, C., Zhao, W., 2022. Transfer learning with artificial neural networks between injection molding processes and different polymer materials. *J. Manuf. Process.* 73, 395–408.
- Malhotra, R., Saxena, I., Ehmann, K., Cao, J., 2013. Laser-induced plasma micro-machining (LIPMM) for enhanced productivity and flexibility in laser-based micro-machining processes. *CIRP Ann.* 62, 211–214.
- Menon, N., Mondal, S., Basak, A., 2022. Multi-fidelity surrogate-based process mapping with uncertainty quantification in laser directed energy deposition. *Materials* 15, 2902.
- Merklein, M., Schulte, R., Papke, T., 2021. An innovative process combination of additive manufacturing and sheet bulk metal forming for manufacturing a functional hybrid part. *J. Mater. Process. Technol.* 291, 117032.
- Misaka, T., Herwan, J., Ryabov, O., Kano, S., Sawada, H., Kasashima, N., Furukawa, Y., 2020. Prediction of surface roughness in CNC turning by model-assisted response surface method. *Precis. Eng.* 62, 196–203.
- Mozaffar, M., Ndiip-Agbor, E., Lin, S., Wagner, G.J., Ehmann, K., Cao, J., 2019. Acceleration strategies for explicit finite element analysis of metal powder-based additive manufacturing processes using graphical processing units. *Comput. Mech.* 64, 879–894.
- Mozaffar, M., Liao, S., Lin, H., Ehmann, K., Cao, J., 2021. Geometry-agnostic data-driven thermal modeling of additive manufacturing processes using graph neural networks. *Addit. Manuf.* 48, 102449.
- Mukherjee, T., DebRoy, T., 2019. A digital twin for rapid qualification of 3D printed metallic components. *Appl. Mater. Today* 14, 59–65.
- Pallav, K., Saxena, I., Ehmann, K.F., 2015. Laser-induced plasma micromachining process: principles and performance. *J. Micro Nano-Manuf.* 3.
- Pan, S.J., Yang, Q., 2010. A survey on transfer learning. *IEEE Trans. Knowl. Data Eng.* 22, 1345–1359.
- Pardoe, D., Stone, P., 2010. Boosting for Regression Transfer, *Proceedings of the Twenty-Seventh International Conference on Machine Learning, ICML 10, Haifa, Israel.*
- Pegues, J.W., Melia, M.A., Puckett, R., Whetten, S.R., Argibay, N., Kustas, A.B., 2021. Exploring additive manufacturing as a high-throughput screening tool for multiphase high entropy alloys. *Addit. Manuf.* 37, 101598.
- Peherstorfer, B., Wilcox, K., Gunzburger, M., 2018. Survey of multifidelity methods in uncertainty propagation, inference, and optimization. *Siam Rev.* 60, 550–591.
- Pokluda, O., Bellehumeur, C.T., Vlachopoulos, J., 1997. Modification of Frenkel's model for sintering. *AIChE J.* 43, 3253–3256.
- Qasimeh, M., Ravoori, D., Jain, A., Adnan, A., 2022. Modeling the effect of in situ nozzle-integrated compression rolling on the void reduction and filaments-filament adhesion in Fused Filament Fabrication (FFF). *Multiscale Sci. Eng.* 4, 37–54.
- Ramezankhani, M., Crawford, B., Narayan, A., Voggenteiter, H., Seethaler, R., Milani, A. S., 2021. Making costly manufacturing smart with transfer learning under limited data: A case study on composites autoclave processing. *J. Manuf. Syst.* 59, 345–354.
- Saha, S.K., Wang, D., Nguyen, V.H., Chang, Y., Oakdale, J.S., Chen, S.-C., 2019. Scalable submicrometer additive manufacturing. *Science* 366, 105–109.
- Saunders, R., Rawlings, A., Birnbaum, A., Iliopoulos, A., Michopoulos, J., Lagoudas, D., Elwany, A., 2022. Additive manufacturing melt pool prediction and classification via multifidelity gaussian process surrogates. *Integr. Mater. Manuf. Innov.* 1–19.
- Saunders, R.N., Teferra, K., Elwany, A., Michopoulos, J.G., Lagoudas, D., 2023. Metal AM process-structure-property relational linkages using Gaussian process surrogates. *Addit. Manuf.* 62, 103398.
- Saxena, I., Ehmann, K., Cao, J., 2014. Laser-induced plasma in aqueous media: numerical simulation and experimental validation of spatial and temporal profiles. *Appl. Opt.* 53, 8283–8294.
- Serdeczny, M.P., Comminal, R., Pedersen, D.B., Spangenberg, J., 2018. Experimental validation of a numerical model for the strand shape in material extrusion additive manufacturing. *Addit. Manuf.* 24, 145–153.
- Wang, J., Zou, B., Liu, M., Li, Y., Ding, H., Xue, K., 2021. Milling force prediction model based on transfer learning and neural network. *J. Intell. Manuf.* 32, 947–956.
- Wang, X., Ma, C., Li, C., Kang, M., Ehmann, K., 2018. Influence of pulse energy on machining characteristics in laser induced plasma micro-machining. *J. Mater. Process. Technol.* 262, 85–94.
- Xie, J., Ehmann, K., Cao, J., 2020. Simulation of ultrashort laser pulse absorption at the water-metal interface in laser-induced plasma micromachining. *J. Micro Nano-Manuf.* 8.
- Xue, T., Gan, Z., Liao, S., Cao, J., 2022. Physics-embedded graph network for accelerating phase-field simulation of microstructure evolution in additive manufacturing. *npj Comput. Mater.* 8, 201.
- Yan, J., Yan, W., Lin, S., Wagner, G.J., 2018. A fully coupled finite element formulation for liquid-solid-gas thermo-fluid flow with melting and solidification. *Comput. Methods Appl. Mech. Eng.* 336, 444–470.
- Yi, W., Jiang, H., Cheng, G.J., 2022. Mesoporous LDH metastructure from multiscale assembly of defective nanodomains by laser shock for oxygen evolution reaction. *Small* 18, 2202403.
- Zhang, Y., Bhandari, S., Xie, J., Zhang, G., Zhang, Z., Ehmann, K., 2021. Investigation on the evolution and distribution of plasma in magnetic field assisted laser-induced plasma micro-machining. *J. Manuf. Process.* 71, 197–211.
- Zhu, Q., Liu, Z., Yan, J., 2021. Machine learning for metal additive manufacturing: predicting temperature and melt pool fluid dynamics using physics-informed neural networks. *Comput. Mech.* 67, 619–635.