

Chapter

Development of a Fast-Solving Machine Learning Surrogate Model for a Pharmaceutical Manufacturing Digital Twin

Mohammad Zandi and Donald Ntamo

Abstract

High-fidelity models (HFM) for twin-screw wet granulation (TSWG) are often too computationally expensive for routine calibration, optimization, and digital twin deployment. This chapter presents a faster, cheaper, and easier-to-use surrogate modeling workflow that preserves HFM-level predictive capability. Sensitivity analysis is applied to prioritize influential inputs, and a rate mechanism global system analysis reduces the calibration space by 60%, leaving only 10 critical parameters. Using this reduced set, a compact mechanistic surrogate comprising 10 equations achieves accurate particle size distribution (PSD) predictions with an R^2 of 0.998 and a simulation time of 0.41 seconds. In addition, a hybrid TSWG model is developed by coupling a population balance model (PBM) with an artificial neural network (ANN). The ANN provides rapid first-pass estimates of kinetic rate parameters within the calibration range, which are then integrated into the PBM to refine PSD predictions and accelerate calibration. The proposed workflow, adaptive framework for robust and accelerated model evaluation (AFRAME), offers a structured and efficient methodology for model calibration in continuous pharmaceutical manufacturing and supports the development of robust, adaptive digital twins applicable to a broader range of manufacturing processes.

Keywords: twin-screw wet granulation, surrogate modeling, population balance model, artificial neural network, digital twin calibration

1. Introduction

The chemical and biochemical industries are undergoing a significant transition, driven by the need to deliver higher-value products with lower environmental impact and greater resource efficiency. Within pharmaceutical manufacturing, this transition has accelerated interest in continuous processing as an alternative to conventional batch production, motivated by the potential for improved throughput, tighter quality control, smaller equipment footprints, and reduced waste [1, 2]. Among the enabling unit operations, twin-screw wet granulation (TSWG)

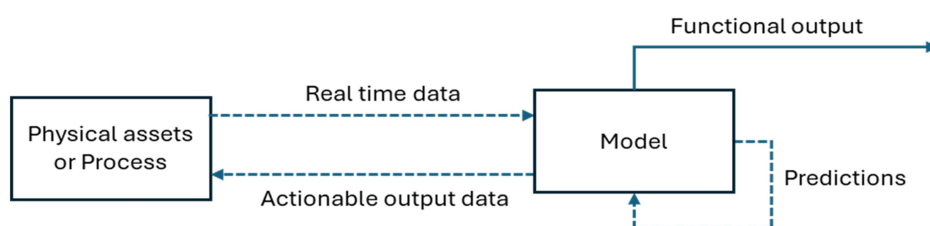


Figure 1.
Predictive digital twin infrastructure.

(see **Figure 1**) has attracted sustained attention because it is well-suited to continuous manufacture and can be integrated with upstream feeding and downstream drying and milling operations [3]. Twin-screw wet granulation offers short residence times, good mixing of active pharmaceutical ingredients (APIs), and enhanced control of granule attributes, making it an attractive route for robust granule production in modern solid-dose platforms [3].

Despite these advantages, the move from batch to continuous manufacturing introduces new technical and operational challenges. Continuous processes are typically more tightly coupled, more sensitive to disturbances, and more susceptible to drift or disruptions than batch operations, particularly when material properties vary or when complex hydrodynamics and kinetics govern product formation. These characteristics increase the burden on process monitoring, control, and rapid decision-making. In parallel, the pharmaceutical sector is experiencing a digital transformation that offers practical mechanisms to address these challenges, including advanced sensing, connectivity, and model-enabled decision support. Digital manufacturing can be understood as the coordinated deployment of information technology (IT) and operational technology (OT), cyber-physical systems, sensing and actuation, connectivity, and data-driven analytics across the value chain to generate value for manufacturing and supply-chain performance in various ways [4]. A key enabler of this digital shift is the digital twin (DT) (see **Figure 1**), which provides a virtual mirror of physical assets to enhance decision-making. In this context, DTs have emerged as a prominent concept for supporting continuous processing and biomanufacturing, enabling faster process development, more systematic optimization, and increased product and supply flexibility.

Digital twins are widely viewed as a key enabler of continuous manufacturing; however, their development can be costly and time-intensive, particularly when rigorous model calibration and validation are required. When time-to-market is a priority, the time needed to build and deploy high-fidelity digital representations becomes a critical constraint.

Models are central to this vision because they support process understanding, enable optimization, and underpin advanced control strategies. However, model choice involves trade-offs. Purely data-driven models (e.g., deep neural networks) can offer strong predictive performance within the domain covered by training data, but they often lack interpretability and may provide limited mechanistic insight into the causal drivers of product quality. Conversely, first-principles (mechanistic) models embed physical and chemical understanding, support extrapolation within reason, and can improve transparency and trust. The limitations of data-driven approaches can, therefore, be mitigated by incorporating first-principles knowledge

whenever feasible, resulting in hybrid strategies that balance interpretability, accuracy, and deployability.

$$\frac{\partial}{\partial t} n_i(v, t) + \frac{\partial}{\partial v} [n_i(v, t) G_i(v)] = \frac{B_{nuc, i}(v, t) + B_{break, i}(v, t) - D_{break, i}(v, t) + n_i(v)}{-n_i - 1(v) \tau_i} \quad (1)$$

where v = Granule volume and t = time:

$\frac{\partial}{\partial t} n_i(v, t)$ = Granule number density of species with volume v at time t .

$G_i(v)$ = Growth rate by layering and consolidation.

$B_{nuc, i}(v, t)$ = Birth rate of granules due to nucleation.

$B_{break, i}(v, t)$ and $D_{break, i}(v, t)$ = Birth and death of granules due to breakage.

τ_i = Mean residence time.

Equation 1 presents a one-dimensional population balance model (PBM). For TSWG, PBMs have been developed and progressively refined to represent granule growth ($G_i(v)$), breakage ($D_{break, i}$), and consolidation mechanisms, with published formulations increasing in sophistication over time [5–8]. Notably, Wang et al. [8] introduced a compartmental approach to capture the influence of screw configuration on granule properties, demonstrating that screw design can materially affect particle size outcomes and should be reflected in the model structure. While incorporating such effects can improve predictive power and enable screw configuration optimization, it typically increases model complexity and calibration burden.

Process systems engineering (PSE) has long provided methods – such as equation-oriented simulation, flexibility analysis, feasibility analysis, and constrained global optimization – to design resilient and cost-effective processes in other sectors, but these approaches have historically seen slower uptake in pharmaceutical manufacturing [9, 10]. Feasibility and flexibility analysis can help identify viable operating regions and design parameters, while optimization algorithms can support systematic recommendations for robust process design [10]. System energy engineering, sensitivity analysis, flow sheeting, and reverse engineering are interconnected disciplines that play crucial roles in optimizing manufacturing systems. Reverse engineering, illustrated in **Figure 2**, is the process of analyzing a system to understand its components and how they interact, often with the goal of improving or replicating it, and will play a crucial role in the future of integrated supply-chain models.

Sensitivity analysis is particularly valuable for product quality management because it identifies which operating variables and kinetic parameters exert the greatest influence on critical quality attributes (CQAs), thereby informing control strategies and experimental planning. For example, design of experiments (DoE) combined with sensitivity analysis has been used to identify critical process parameters (CPPs) in twin-screw granulation, and global sensitivity analysis (GSA) using Sobol' indices has been applied to improve the efficiency of model-based design. In their wet granulation case study, inputs such as flow rate, liquid-to-solid ratio, screw speed, dryer air temperature, and drying time were evaluated, with the variance-based analysis highlighting the liquid-to-solid ratio as a dominant factor [11].

Alongside mechanistic modeling, artificial neural networks (ANNs) have become widely used in process modeling due to their ability to approximate complex nonlinear relationships. Artificial neural networks typically consist of an input layer, one or more hidden layers, and an output layer, with model parameters (weights)

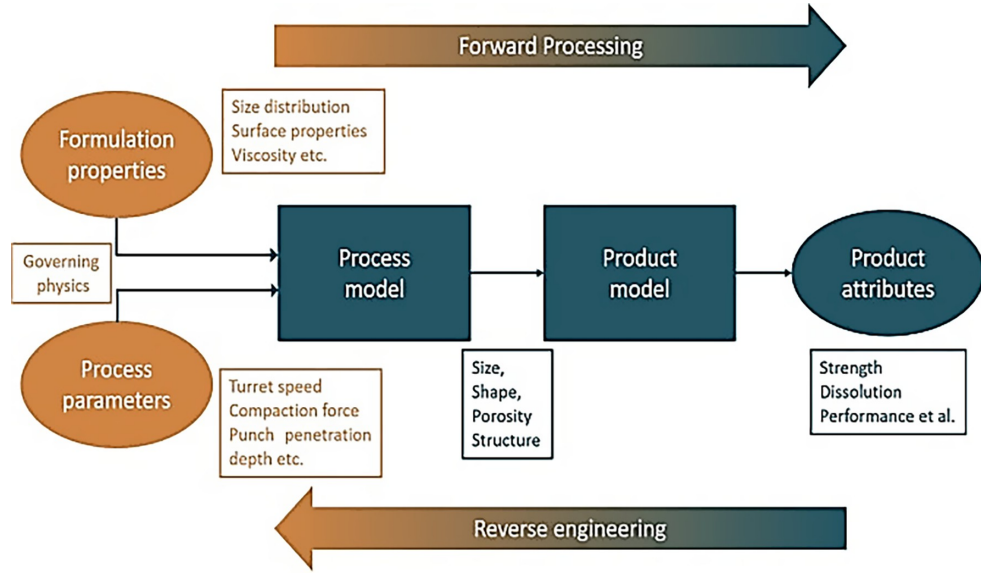


Figure 2. Formulated product reverse engineering: a process of inferring formulation properties from product attributes. Source: Study by Litster and Bogle [9].

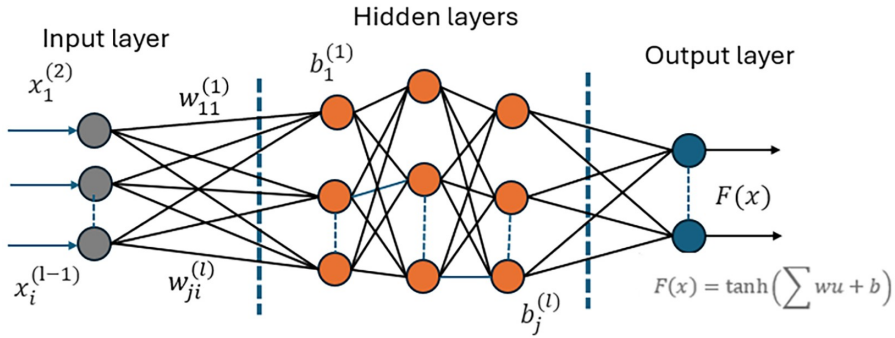


Figure 3. Machine learning algorithm structure: time t included for dynamic modeling and excluded for steady-state modeling.

adjusted during training to map inputs to outputs, as illustrated in **Figure 3**. A well-known limitation is overfitting, where the network fits noise in the training set and generalizes poorly. Techniques such as Bayesian regularization have been developed to reduce overfitting and improve generalization performance [12]. In addition, modern open-source ecosystems and libraries (e.g., TensorFlow and Keras) have lowered the barrier to implementing and training neural networks using widely adopted computing environments such as Python and R.

Although rigorous mechanistic models can now be constructed for many unit operations, two recurring issues arise when these models are intended for online digital operations. First, high-fidelity models (HFMs) can be computationally

expensive, making them unsuitable for use cases that require rapid simulation, repeated optimization, or real-time decision-making.

Second, the standard model development workflow for complex models, as illustrated in **Figure 4**, may lack numerical robustness across the full operating envelope, particularly when calibration parameters are uncertain or when the model must handle noisy, time-varying industrial data. These limitations are especially consequential for model predictive control (MPC), which relies on repeatedly forecasting process behavior within tight time constraints. If a model cannot deliver reliable predictions within seconds, it becomes impractical for industrial MPC deployment, thereby motivating the development of faster model alternatives.

Surrogate modeling methodology, illustrated in **Figure 5**, provides a pragmatic route to address these constraints by replacing computationally intensive mechanistic simulators with computationally efficient approximations – often machine-learning-based – trained to reproduce HFM outputs with quantifiable accuracy over a defined domain. When designed carefully, surrogate models can reduce simulation times by orders of magnitude, enabling rapid optimization, scenario evaluation, and control-oriented computations. In the context of DTs, surrogates can also support earlier integration into digital architectures, accelerating the path from process development to operational deployment.

Against this backdrop, this chapter focuses on accelerating the digital deployment of high-fidelity models for twin-screw wet granulation by developing a surrogate modeling methodology suitable for continuous pharmaceutical manufacturing. The work first applies global sensitivity analysis as a core element of model development and calibration, particularly relevant when many inputs and kinetic parameters may influence outputs. Building on this, machine learning (ML) techniques are used to construct computationally efficient surrogate models that reproduce particle size distribution predictions with high accuracy while substantially reducing computational cost. The chapter further introduces a hybrid modeling strategy that couples PBM structure with ANN capability, using ML to provide rapid initial estimates of kinetic parameters that can be embedded within a mechanistic PBM to streamline calibration. Finally, the chapter presents a workflow for integrating rigorous mechanistic models into digital architectures through surrogate modeling, and it discusses practical limitations and future research directions for scalable, robust DTs in pharmaceutical and broader chemical manufacturing applications.

2. Experimental setup

The Diamond Pilot Plant (DiPP) at the University of Sheffield houses large-scale equipment serving multiple industrial sectors, with a particular emphasis on pharmaceutical manufacturing [14]. The facility includes a continuous crystallization unit, a continuous filter dryer, and an industrial-scale GEA ConsiGma™ 25 powder-to-tablet line. The continuous oscillatory baffled crystallizer (COBC) and the continuous filter dryer enable the production of high-purity crystals and drug substances for downstream drug product manufacture. The DiPP continuous powder processing plant demonstrates the growing importance of complex particulate and formulated products within modern chemical engineering.

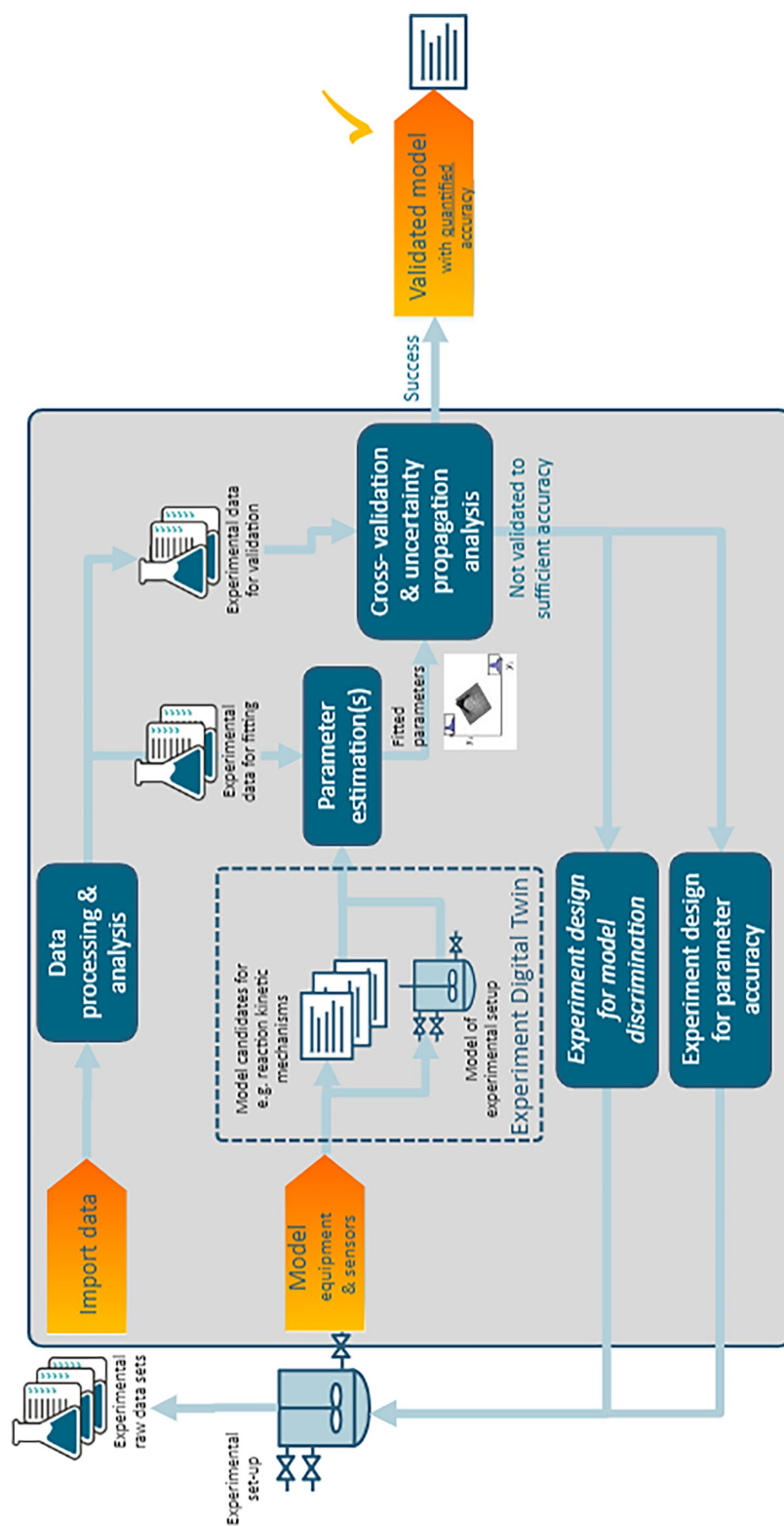


Figure 4.
Iterative model development and validation procedure. Source: Siemens Digital Industries Software [13].

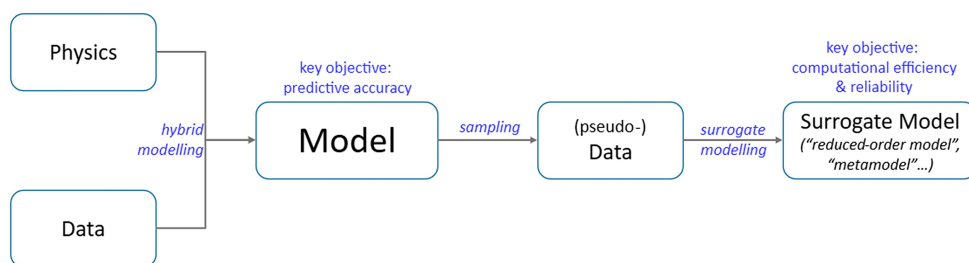


Figure 5.
 Surrogate modeling workflow.

It is supported by an advanced monitoring and control infrastructure, including a dedicated industrial control room that enables state-of-the-art advanced process control. Collectively, these capabilities position the DiPP as a leading Industry 4.0 demonstrator [14]. A key objective of the platform is to implement a data-driven and automated framework that allows prospective Industrial internet of things (IIoT) applications to be tested and validated under realistic operating conditions. Granulation is a particle size-enlargement process in which fine powders are transformed into larger agglomerates (granules). In the manufacture of solid oral dosage forms, granulation is commonly applied to improve powder flowability, enhance compaction behavior, and promote blend uniformity for downstream processing. Granulation may be performed as either a dry or a wet process.

In wet granulation (see **Figure 6**), granules form primarily through nucleation, initiated by the addition of a liquid binder onto an agitated powder bed. In a twin-screw granulator, powder and liquid are mixed under co-current screw conveyance using a sequence of conveying and kneading elements (KEs). The combined effects of wetting and mechanical shear promote particle rearrangement and growth mechanisms – such as layering and consolidation – while simultaneous breakage and re-agglomeration help establish the evolving granule size distribution. The outcome is a population of wet granules with properties suitable for subsequent drying, milling, and tableting operations.

A 20-liter tumble blender (Inversina-Bioengineering, Wald, Switzerland) was used to prepare the formulation preblend. The blender was operated at 25 rpm for 15 minutes, after which the blend was transferred to a gravimetric loss-in-weight feeder (KT20, K-Tron Soder, Niederlenz, Switzerland). The powder formulation comprised 72% lactose (α -lactose monohydrate, Pharmatose 200 M), 24% microcrystalline cellulose (MCC; Chemical PH 101), and 4% polyvinylpyrrolidone (PVP; Povidone K-25). As this work focused exclusively on granulation, the twin-screw granulator (TSG) was operated as a standalone unit and disconnected from the remainder of the ConsiGma™ line. Demineralized water was used as the granulation liquid and was gravimetrically dosed into the granulator using two out-of-phase peristaltic pumps (Watson-Marlow, Cornwall, UK). The screw configuration comprised eight compartments: five containing conveying elements (CEs) and three containing KEs, as shown in **Figure 7**.

The three parameters (powder feed rate, L/S ratio, and screw speed) were chosen because research supports their strong influence on torque and PSD [11]. See **Table 1**.

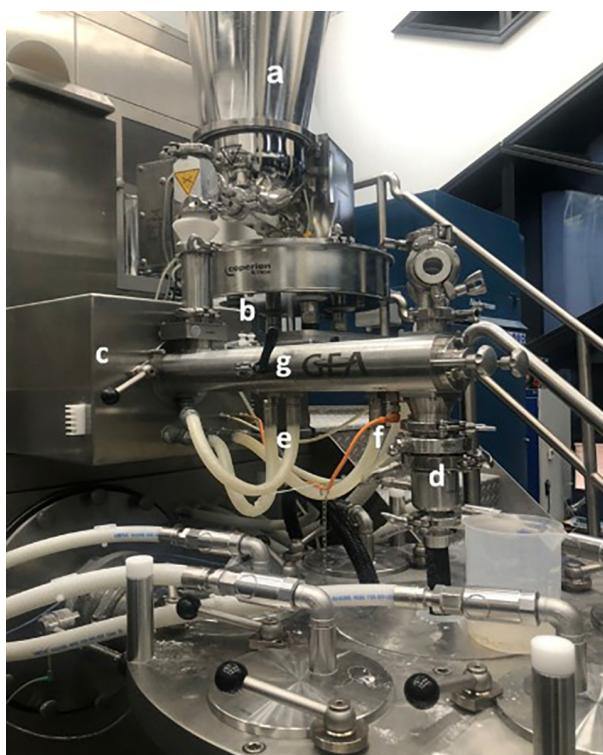


Figure 6.
Diagram of a typical TSWG unit. (a) Powder feed; (b) liquid feed; (c) motor unit; (d) granule outlet; (e) temperature control inlet; (f) temperature control outlet; and (g) TSWG barrel.

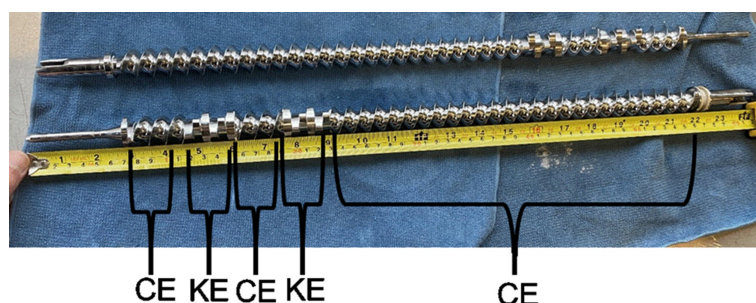


Figure 7.
Configuration and sizes of kneading elements (KE) and conveying elements (CE) used in the twin-screw granulator.

The maximum liquid-to-solid ratio (L/S) to make granules and prevent paste is 0.45. Low screw speed combined with a high powder flow rate results in higher torque, and it may exceed the safety limit of the equipment; therefore, the machine may stop, and you should avoid these conditions.

All kneading blocks were installed at a forward stagger angle of 60° and were separated by CEs. The granulator barrel jacket was preheated and maintained at

Parameter	Range
Screw speed (rpm)	200 to 900
Powder feed rate (kg/hr)	5 to 25
Liquid feed rate (g/min)	15 to 208

Table 1.
Twin-screw wet operational boundaries.

25°C. Unless otherwise stated, operating conditions were held constant at a screw speed of 500 rpm and a fixed powder mass feed rate during each run. To explore process sensitivity, the liquid-to-solid (L/S) ratio, powder feed rate, and screw speed were varied over the ranges of 0.15–0.35, 5–20 kg/h, and 200–500 rpm, respectively.

Prior to granulation, the particle size distribution (PSD) of the formulation preblend was characterized by sieve analysis to provide baseline material properties; these results were subsequently used as inputs to gPROMS FormulatedProducts. Although image-based systems such as Computer-Aided Material Sizer (CAMSIZER) provide higher-resolution, automated PSD and shape information, sieve analysis remains a robust, cost-effective, and widely standardized method for routine PSD measurement when detailed morphology is not required.

Wet granules were collected at the TSG outlet and dried overnight at ambient DiPP conditions (22 °C) prior to characterization. Granule size distributions were measured using a sieve shaker with a stack of sieves spanning 60–2,800 µm. The sieve stack was shaken at an amplitude of 0.76 mm with a 20-second interval setting for a total duration of 1 minute. The retained mass on each sieve was weighed, and the granule size distribution was reported as mass fraction versus sieve size.

3. High-fidelity modeling

Siemens’ gPROMS FormulatedProducts has been widely adopted for simulation and digital applications in the pharmaceutical sector and is frequently cited in the modeling literature. Representative examples include work on modeling particulate processes for the continuous manufacture of solid dosage forms, surrogate modeling of dissolution behavior to support efficient tablet process design, and Pharma 4.0 DT concepts for solidified nanosuspensions [15, 16] (Rogers et al., 2013 [17]). gPROMS is used by a broad industrial and academic community, including Fortune 500 process-industry companies and approximately 200 universities, supported through operations in the UK, USA, UAE, Japan, and Korea, with additional agency coverage in China and Taiwan (PSE: Sanofi Joins Systems-based Pharmaceuticals Alliance, 2019 [18]). This established deployment base, alongside substantial industrial implementation experience, underpins its suitability as a platform for DT development.

From a compliance perspective, Siemens’ solutions are positioned to support regulated environments, with alignment to key requirements for computerized systems used in pharmaceutical manufacturing, notably EU GMP Annex 11 (“Computerized Systems”) and the US FDA’s 21 CFR Part 11 (“Electronic Records; Electronic Signatures”) (Good Manufacturing Practice, Siemens.com) [19, 20]. These frameworks define acceptance criteria for validated computer systems,

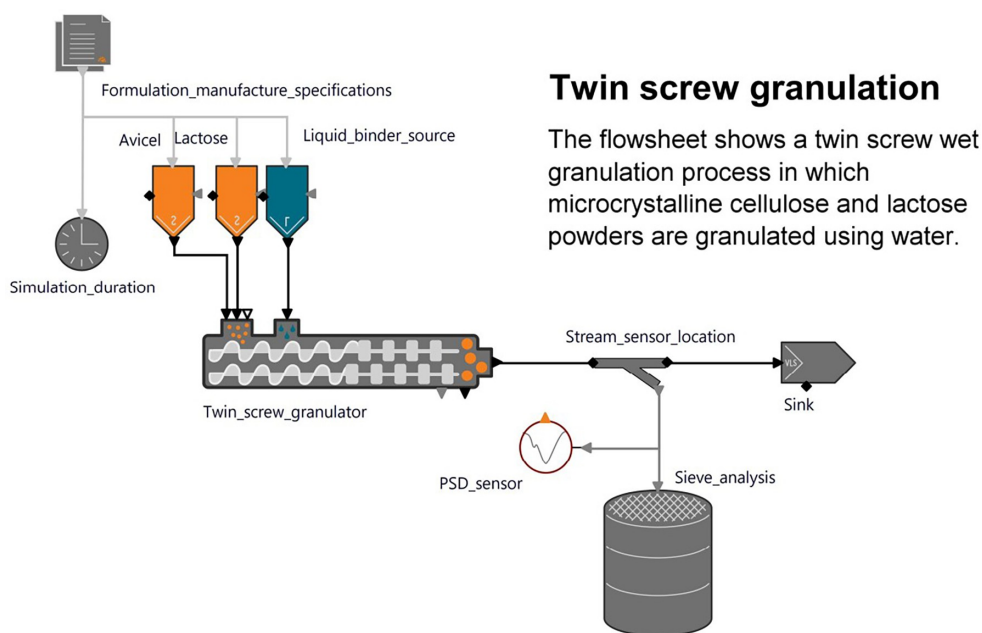


Figure 8.
Flowsheet model of twin-screw granulator on Siemens' gPROMS FormulatedProducts.

electronic records, and electronic signatures in GxP contexts. gPROMS FormulatedProducts provides a library of validated modeling components to support the simulation and optimization of pharmaceutical unit operations, including fluidized-bed dryers, mills, and granulators such as TSWGs. Modeling within gPROMS also enables the integration of digital MPC capabilities through linkage with the online PharmaMV suite, supporting real-time decision-making and control deployment. In addition, the platform includes optimization functionality for defining and minimizing objective functions (e.g., deviation from target CQAs or operating cost proxies), facilitating systematic process design and improvement.

Siemens' gPROMS FormulatedProducts was used to develop a HFM of the DiPP ConsiGma™ 25 TSW unit operation, enabling the prediction of the PSD, as illustrated in **Figure 8**. The platform provides advanced parameter estimation capabilities to fit kinetic parameters with high precision and includes a built-in global analysis module to support sensitivity studies and the systematic assessment of input–output relationships.

4. ML and surrogate modeling

The primary objective of the surrogate modeling approach is to retain the predictive accuracy of the HFM while substantially reducing computational cost. Using appropriate sampling strategies, the calibrated mechanistic model can be interrogated to generate a sufficiently rich set of synthetic (pseudo) input–output data. These datasets are then used to train a computationally efficient and, where possible, more numerically robust surrogate model. In practice, surrogate models can

be solved considerably faster than the underlying mechanistic model; however, their accuracy depends on (i) the coverage and quality of the sampling used to generate training data and (ii) the choice and configuration of the data-driven modeling method.

The workflow for developing the surrogate model flowsheet can be summarized in four stages:

Stage 1: Develop a working flowsheet or custom model grounded primarily in first-principles.

Stage 2: Use global system analysis to generate extensive synthetic input–output datasets from the mechanistic model.

Stage 3: Train and validate candidate surrogate models using the generated datasets.

Stage 4: Deploy the selected surrogate within the flowsheet to enable rapid simulation, optimization, and DT integration.

Neural networks can be thought of as layers of computation units, neurons with connections between consecutive layers. The network is comprised of one input layer with a predetermined number of neurons, which is equal to the number of model input parameters, one or more hidden layers with a custom number of neurons, and one output layer.

$$z_j^{(l)} = \sum_{i=1}^{N_l} w_{ji}^{(l)} x_i^{(l-1)} + b_j^{(l)} \text{ for } j = 1, \dots, J_l \quad (2)$$

where $z_j^{(l)}$ is activation j at layer l , $x_i^{(l-1)}$ is the input from layer $l - 1$ to layer l , $w_{ji}^{(l)}$ is a weight parameter at layer l , and $b_j^{(l)}$ is a bias parameter.

In the network configuration, the user can see the input and output layers, which are automatically generated by gPROMS, and can add one or more hidden layers. gPROMS currently supports numerous types of ANN activation functions, either linear or nonlinear. The configuration of ANNs allows for significant flexibility in how data is processed through the hidden layers. Activation functions are the mathematical gates that determine whether a neuron should be activated based on whether the input it receives is relevant to the model's prediction. The Rectified Linear Unit, or ReLU, is currently the most widely used activation function in deep learning. It outputs the input directly if it is positive; otherwise, it outputs zero. In gPROMS, this function is highly valued for its computational efficiency, as it allows models to train faster by avoiding complex exponential calculations. Because it “turns off” neurons that produce negative values, it creates a sparse network that is less prone to the “vanishing gradient” problem often found in deeper architectures. The Sigmoid activation function, also known as the logistic function, takes any real-valued input and squashes it into a range between 0 and 1. This makes it particularly useful for models where the output needs to be interpreted as a probability. In the context of gPROMS simulations, Sigmoid is often used in the output layers of binary classification tasks or within hidden layers when the user wants to ensure that the signals remain within a strictly bounded, positive range. However, because the function flattens out at the extremes, it can sometimes lead to slower learning during the backpropagation process. The hyperbolic tangent, or Tanh, function is similar in shape to the Sigmoid but maps input values to a range between -1 and 1 . This zero-centered nature is a distinct advantage in many engineering applications within gPROMS, as it ensures that the mean of the activations remains near zero, which generally helps the model converge more quickly during training. By allowing for negative outputs, the Tanh function can more effectively capture relationships in

data where the direction of the signal (positive vs. negative) is just as important as its magnitude.

A crucial step in developing any ML model is evaluating its accuracy. The model's predictive capability is assessed by comparing its predictions to observed or validated synthetic values. Common metrics used for this evaluation include mean absolute error (MAE), relative absolute error (RAE), root-mean-square error (RMSE), and R-squared (R^2).

$$MSE = \frac{\sum_i |f_i - y_i|}{n} \quad (3)$$

$$RAE = \frac{\sqrt{\sum_i (f_i - y_i)^2}}{\sqrt{\sum_i (y_i)^2}} \quad (4)$$

$$RMSE = \sqrt{\frac{\sum_i (f_i - y_i)^2}{n}} \quad (5)$$

$$R^2 = 1 - \frac{\sum_i (f_i - y_i)^2}{\sum_i (f_i - \bar{y})^2} \quad (6)$$

where f and y refer to predicted and observed values, respectively, i also refers to the set of experimental runs, and n refers to the number of measurements.

4.1 TSWG mechanistic model

Given the screw configuration's significant impact on granule size distribution, a compartmental approach is typically used to evaluate material transport along the barrel. Each axial compartment is assigned an average residence time based on the screw elements. The inlet flow rate in one compartment is equal to the outlet flow rate in the last neighboring compartment, as described by Eq. (6).

$$\dot{M}_{out} = \frac{M}{\tau} \quad (7)$$

where $\dot{M}_{out}$ is the outlet flow rate for each compartment; M is the mass holdup, and τ is the residence time in each compartment. The inlet flow rate of a compartment is equal to the outlet flowrate of the previous one. However, there is an exception for the first compartment where the inlet flow rate is specified based on the inlet feed rate of materials. The following models, selected to describe the PBM rate mechanism, are grounded in the fundamental processes of nucleation, breakage, layering, and consolidation. These models represent advancements over previous iterations, incorporating updated mechanisms to provide a more comprehensive understanding of the PBM process.

Drop nucleation, adapted from Barrasso and Ramachandran [5], involves the formation of droplets with a lognormal size distribution, which penetrate a bed of fine powder and replace interparticle voids. The probability density function is given as follows:

$$f(\mu_{\text{drop}}) = \frac{1}{\mu_{\text{drop}}} \frac{1}{\sigma_{\text{drop}} \sqrt{2\pi}} \exp\left(-\frac{(\ln \mu_{\text{drop}} - \mu_{\text{drop, mean}})^2}{2\sigma_{\text{drop}}^2}\right) \quad (8)$$

where $\mu_{\text{drop, mean}}$ is the mean droplet size and σ_{drop} is the standard deviation. The rate of nuclei formation in compartment i , $B_{\text{nuc},i}$ is determined by droplet size μ_{drop} , bed porosity ϵ_{bed} , and maximum pore saturation using Eq. (8) by assuming that the rate of nuclei formation is equal to the rate of droplet addition to the powder fines [5].

$$B_{\text{nuc},i} = \frac{\dot{m}l_i}{V_{\text{drop}}} \delta(v - v_{\text{nuc}}) \quad (9)$$

where l_i is the flow rate of liquid added to compartment i , m is the mass in the compartment, V_{drop} is the volume of a single liquid drop, v is the granule volume, and v_{nuc} is the saturated granule volume. The nuclei volume is determined from Eq. (10), which assumes that each nucleus is fully saturated with liquid and has a porosity equal to that of the bed.

$$v_{\text{nuc}} = \frac{V_{\text{drop}}}{s * \epsilon_{\text{bed}}} \quad (10)$$

Where s^* is the droplet pore penetration, and ϵ_{bed} is the bed porosity. Granule breakage, modeled after Wang et al. (2020) [21], considers chipping and fragmentation mechanisms in conveying and distributive mixing elements. A Weibull distribution is used to describe the cumulative breakage function $P(b, \text{KE})$ for the distributive KE due to its versatility in modeling various failure or breakage behaviors, while a bimodal distribution represents the coarse and fine particles resulting from chipping in the CE (Wang et al., 2020 [21]).

The breakage probability $P(b, \text{CE})$ was found to increase with an increase in powder feed number (PFN) and granule size and a decrease in the dynamic yield stress of the granule (DYS). Taking these parameters into account, along with the upper and lower granule size limits for breakage, x_{lc} and x_{uc} , an empirically derived selection function for CE may be written as Eq. (11).

$$P_{b, \text{CE}} = \begin{cases} 1 & x \geq x_{uc} \\ a \times \exp\left(-\frac{\text{DYS}}{b} \times \frac{1}{\text{PFN}}\right) \times (x/x_{lc})^c & x_{lc} < x < x_{uc} \\ 0 & x \leq x_{lc} \end{cases} \quad (11)$$

where a , b , and c are fitting parameters in the breakage kernel and represent the lower and upper granule size bounds. To represent the distribution of coarse and fine particles that result from chipping in CEs, a bimodal distribution is proposed:

$$B_{M, \text{CE}} = vB_{M1} + (1 - v)B_{M2} \quad (12)$$

where v is the proportion of fine and coarse particles, and $BM1$ and $BM2$ are the modified Weibull distributions.

The breakage probability in the KE of the screw, $P(b, \text{KE})$, is found to be a strong function of the screw geometry, DYS , and the PFN. We propose the following empirically derived three-parameter model for the breakage selection function.

$$P_{b, CE} = \begin{cases} 1 & x \geq x_{uc} \\ 1 - a' \times \exp\left(-PFN \times \frac{b'}{DYS} \times \left(\frac{x}{x_{lc}}\right)^{c'}\right) & x_{lc} < x < x_{uc} \\ 0 & x \leq x_{lc} \end{cases} \quad (13)$$

A modified Weibull distribution is proposed to describe the breakage fragment size distribution via fragmentation and is represented by Eq. (14).

$$B_{M, KE} = \exp\left(-\beta \times \left(\frac{x}{x_{uc}}\right)^\gamma\right) \quad (14)$$

where BM,KE is the cumulative size distribution, β is the scale parameter, and γ is the shape parameter in the Weibull function.

Layering, based on 6, involves the deposition of fine powder particles onto wet granule surfaces, leading to an increase in granule size, as described by Eq. (15):

$$G_i(x) = G_m \frac{M_{i,p}}{kM_{i,g} + M_{i,p}} \exp[-\lambda(y_{w,i} - y_{wc})^2] \quad (15)$$

where $G_i(x)$ is the linear size-based growth rate, G_m is the layering maximum linear growth rate, λ and k are fitting parameters in the layering kernel, $M_{i,g}$ is the mass of granules in compartment i , $M_{i,p}$ is the mass of fine powder in compartment i , $y_{w,i}$ is the moisture content, and y_{wc} is the critical moisture content. Granule growth rate by layering, $G_i(v)$, is related to the linear size-based growth rate using Eq.(16).

$$G_i(v) = 6.9v^{\frac{2}{3}}G_i(x) \quad (16)$$

The consolidation model, adapted from [7], accounts for the reduction in particle size due to decreasing internal porosity. This model considers the combined effect of intra-granule voids and internal liquid, leading to a decrease in liquid and internal pore space. The mass of other granulate components remains conserved, and the loss of liquid from the granulate phase is balanced by an increase in external liquid [7]. The consolidation rate is characterized by an exponential decrease in porosity ϵ_i , approaching a minimum value ϵ_{min} [6] and described by Eq. (17).

$$R_i(\epsilon) = -k_{i, comp}(\epsilon_i - \epsilon_{min}) \quad (17)$$

where $R_i(\epsilon)$ is the rate of change of granule porosity in compartment i , and $k_{i, comp}$ is the compaction rate constant s^{-1}

A total of 25 parameters listed in **Table 2** are to be calibrated. These parameters fall into two main categories: material properties and equipment geometry. Material properties, such as the dynamic yield stress, influence the behavior of the material during the granulation process. Equipment geometry, such as the screw design, affects the breakage and consolidation of granules.

4.2 Rate mechanism global system analysis

The following section summarizes the sensitivity analysis of each rate mechanism based on the defined process kernels in PBM, while holding all other rate mechanisms constant. Lower and upper bounds are defined by shifting each

Rate mechanism	Parameters to be estimated	Symbol	Units
Consolidation	Consolidation minimum porosity	ε_{min}	-
	Consolidation rate constant	k_i	s^{-1}
Layering	Layering critical moisture content	y_{wc}	kg/kg^{-1}
	Layering kinetic parameter λ	λ	-
	Layering kinetic parameter k	k	-
	Layering maximum linear growth rate	G_m	$\mu m s^{-1}$
Nucleation	Drop nucleation means initial droplet size	$\mu_{drop, mean}$	mm
	Drop nucleation pore penetration	v_{nuc}	$m^{-3}m^{-3}$
	Drop nucleation standard of initial droplet size	σ_{drop}	mm
Breakage	Breakage rate constant in CE	$B_{M, CE}$	s^{-1}
	Breakage rate constant in KE	$B_{M, KE}$	s^{-1}
	Breakage dynamic yield strength (DYS) in CE	DYS	kPa
	DYS in KE	DYS	kPa
	Breakage fitting parameter "a" in KE	a'	-
	Breakage fitting parameter c in KE	c'	-
	Breakage fitting parameter c in CE	c	-
	Maximum breakage granule size in KE	x_{lc}	μm
	Maximum breakage granule size in CE	x_{lc}	μm
	Minimum breakage granule size in KE	x_{uc}	μm
	Minimum breakage granule size in CE	x_{uc}	μm
	Breakage shape parameter in KE	γ	-
	Breakage scale parameter in KE	β	-
	Powder Feed Number (PFN) in KE	PFN	-
	PFN in CE	PFN	-
	Fraction of fines in CE	v	-

Table 2.
Model calibration parameters.

parameter by 20% of its initial reference value taken from the literature [8]. The process parameters were specified within the equipment's operational limits, that is, screw speed operating conditions of 200–900 rpm, powder feed rate operating conditions of 5–25 kg/hr, and Liquid to-solid ratio operating conditions of 0–0.45. The DEM-informed parameters were adjusted to cover a wide range of values extracted from various DEM simulations performed in the literature [8].

Variance-based sensitivity indices are used to measure the influence of individual factors on responses. The variance-based sensitivity analysis method is based on Saltelli's method and the formulae used to estimate the sensitivity indices [19]. The indices were computed using the Monte Carlo method. The main drawback of variance-based sensitivity analysis is its high computational cost. For this case study, a quasi-random sampling method that covers the whole operational space with over 50 to 2,000 samples was used. Estimating the sensitivity indices requires many

model evaluations. The main assumption when carrying out the analysis is that rate mechanisms are additive, and the factors are each interactive but not correlated.

The first-order effect, that is, the direct effect of factor i on response y , is given by Eq. (18).

$$S_i = \frac{\sigma_{i,y}^2}{\sigma_y^2} \quad (18)$$

The total effect, that is, the total effect of factor i on the response y , is given by Eq. (19).

$$S_{i,T} = \frac{\sigma_{i,T,y}^2}{\sigma_y^2} \quad (19)$$

The first-order effect index represents the main effect contribution of each input factor to the variance of the output. The higher the value of S_i , the higher the influence of the i th factor on the output. In other words, a high value of S_i indicates an important factor. If $S_i = 0$, then the i th factor has no direct influence on the output; however, a small value of S_i does not necessarily imply a non-influential factor. It may still be an important factor through its interactions with other factors. For this reason, its total effect must also be examined. The total effect index, $S_{i,T}$, accounts for the total contribution to the output variance of the i th factor, including the first-order effect plus all higher-order effects due to its interactions with other factors. If $S_{i,T}$ is equal to S_i , then the factor has no interactions with the other factors. If $S_{i,T} = 0$, the i th factor has no influence on the model output, and the factor can be fixed at any value within its range of uncertainty. For each rate mechanism parameter, the total sensitivity index value per output shows which are the most impactful modeling parameters to be included in parameter estimation. For an impactful parameter, the criteria state that the average total sensitivity index value must be greater than 0.1. The average total sensitivity index value is an average of the D10 total sensitivity index, D50 total sensitivity index, and D90 total sensitivity index. Typically, there are a minimum number of samples required to produce accurate results. The optimal number of samples is case-specific and depends on system nonlinearity, system complexity, factors, and responses selected, as well as their number and method. To check for convergence, the analysis is executed for different numbers of samples by plotting the key metric of interest, the average total effect index, against the number of samples for each rate mechanism. In the first instance, the sample size will be increased from 50 to 2000 samples until the sensitivity index converges (see **Table 3**).

4.3 Surrogate modeling

Synthetic data generation provides a practical and effective means of addressing common ML challenges, particularly when high-quality operational datasets are limited (Klein and Bergmann, 2018 [22]). In this approach, validated process simulators are used to generate realistic yet controlled datasets that systematically span the operating domain. These simulated data can then be combined with plant or pilot-scale measurements to enrich the training set and improve coverage of conditions that are difficult to observe experimentally.

High impact parameters	Average total sensitivity index
Breakage rate constant (CE)	1.4423(1)
Breakage rate constant (KE)	0.0298(7)
Fraction of fine particles (CE)	1.3442(2)
Breakage scale parameter (KE)	0.0009(10)
Breakage shape parameter (KE)	0.0031(9)
Drop nucleation pore penetration	0.7595(5)
Drop nucleation standard deviation of initial droplet size	0.0236(8)
Layering critical moisture content	2.8978(4)
Layering kinetic parameter A	0.3144(6)
Layering maximum growth rate	1.1758(3)

Table 3.
Summary of the influence of the TSG kernel on the output PSD.

Synthetic data are especially valuable when collecting real-world observations is difficult, time-consuming, or cost-prohibitive and when rare or safety-critical events must be represented without exposing equipment or personnel to unnecessary risk (Saxena et al., 2008 [23]). In addition to expanding the range of operating conditions captured in the training set, synthetic data can mitigate class imbalance in classification tasks by augmenting underrepresented regimes, thereby improving model robustness and generalizability.

Following preprocessing and dataset partitioning (training, validation, and testing), the resulting dataset is suitable for training the surrogate model. The training data enable the model to learn the mapping between selected input features (operating conditions and model parameters) and the target outputs (e.g., PSD metrics and other CQAs).

In this study, the surrogate training data were obtained from a global sensitivity analysis performed using a previously calibrated and validated mechanistic model. Simulations were conducted using the optimized parameter values reported in **Table 5** and within the operational boundaries defined in **Table 4**.

4.4 Surrogate model configuration

To develop the ML surrogate, a structured workflow was followed. First, the operating variables and model parameters that most strongly influence twin-screw granulation performance were identified and incorporated as model inputs. **Table 5** summarizes the selected variables, their process relevance, and the manner in which they were represented in the dataset.

4.5 Model evaluation and discussion

Table 6 provides a detailed evaluation of each set: training and validation. This evaluation includes metrics such as RMSE and R^2 , allowing for a clear comparison of their performance. By analyzing these results, one can gain valuable insights into the

Mechanism	Parameters	Optimized value
Nucleation	Droplet mean size standard deviation	0.52
	Droplet pore penetration	0.40
Breakage	Breakage rate in CE	1.5
	Breakage rate in KE	2.5
	Scale parameter in CE	2.54e ¹⁰
	Shape parameter in KE	8.23
	Fraction of fines in CE	0.4
Layering	Growth rate	2666.5

Source: Wang et al. [24].

Table 4.
Twin-screw wet optimized parameters.

Surrogate model inputs		
Symbols	Variables	
U1	Liquid mass flow rate	kg/h
Surrogate model output		
Y1	D10	μm
Y2	D50	μm
Y3	D90	μm
Y4	Mass fraction screen 1(2000 μm)	kg/kg
Y5	Mass fraction screen 1(1,180 μm)	kg/kg
Y6	Mass fraction screen 1(850 μm)	kg/kg
Y7	Mass fraction screen 1(500 μm)	kg/kg
Y8	Mass fraction screen 1(250 μm)	kg/kg
Y9	Mass fraction screen 1(90 μm)	kg/kg
Y10	Mass fraction screen 1(45 μm)	kg/kg
Scaling		
Input scaling		Normalized
Output scaling		Normalized
Number of points		
Used for training		100
Used for cross-validation		200
Neural Network		
Number of hidden layers		2
Layer	Activation function	Number of neurons
Hidden layer 1	Tanh	30
Hidden layer 2	Tanh	30
Output	Linear	1

Table 5.
Surrogate model configuration.

Accuracy metrics	Dataset	MAE	RAE	RMSE	R ²
D10	Training	0.0181078	0.000530503	0.0226737	0.998985
	Cross-validation	0.0184147	0.000538770	0.0230374	0.998973
D50	Training	0.57	0.00442550	0.679053	0.951188
	Cross-validation	0.57	0.00449420	0.689237	0.948575
D90	Training	0.242677	0.000283546	0.280787	0.999966
	Cross-validation	0.244845	0.000286712	0.283553	0.999966

Table 6.
PSD distribution machine learning evaluation table.

strengths and weaknesses of each set, influencing performance areas for improvement.

A critical observation is the stability of the metrics during cross-validation. The MAE and RMSE remained consistent with the training results (e.g., D10 training RMSE of 0.0227 vs. cross-validation RMSE of 0.0230). This parity indicates that the preprocessing steps, such as normalization and feature engineering, effectively prevented overfitting, allowing the model to maintain high numerical stability. Once the ML model has demonstrated robust performance in controlled environments, the next crucial step is to deploy it into real-world applications. This process, known as operationalization, involves integrating the model into existing systems and workflows to deliver practical value.

The predictive capabilities of the deployed models were evaluated, and the results are illustrated in **Figures 9** and **10**.

The validation results demonstrate that the ML surrogate model is a highly effective proxy for the complex HFM in pharmaceutical twin-screw granulation. Specifically, the surrogate model maintains strong predictive accuracy across the PSD, showing a close match to the high-fidelity outputs for the D10, D50, and D90 quartiles in **Figures 9** and **10**. Beyond accuracy, the surrogate model offers a profound leap in computational efficiency by reducing the underlying mathematical complexity from 54,271 equations to a mere 10 (see **Table 7**). This reduction translates to a simulation time of only 0.40672 seconds, compared to the 83.7564 seconds required by the high-fidelity version. This approximately 200-fold increase in speed is the critical factor that enables the DT to function in real-time, allowing for the immediate predictive and interactive feedback loops necessary for modern manufacturing environments.

In this specific case, the developed ML models were seamlessly integrated into the gPROMS process flowsheet as custom models. This integration enabled the models to provide real-time predictions and insights, directly influencing decision-making and optimizing process operations.

By operationalizing the ML models, significant benefits were realized, including:

- Enhanced decision-making: Real-time predictions empowered operators and engineers to make informed decisions, leading to improved process efficiency and product quality.

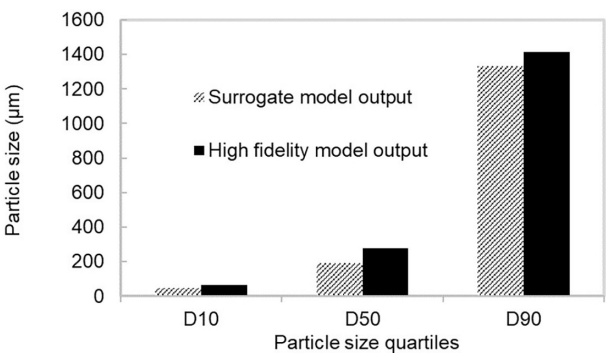


Figure 9.
PSD quartile prediction: surrogate model vs. high-fidelity model.

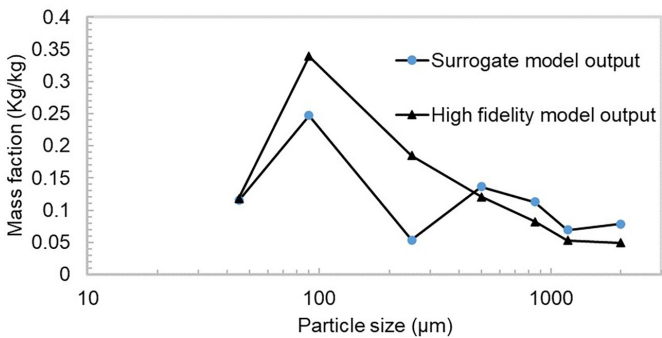


Figure 10.
PSD prediction: surrogate model vs. high-fidelity model.

Model output	Surrogate	High-fidelity model
Number of equations	10	54,271
Simulation time (s)	0.40672	83.7564

Table 7.
Simulation stats: ML surrogate model prediction and high-fidelity model prediction.

- Optimized process control: The models enabled precise control of process variables, minimizing deviations, and maximizing yields.
- Predictive maintenance: By anticipating potential equipment failures, maintenance costs were reduced and downtime was minimized.
- Accelerated product development: The models facilitated rapid experimentation and optimization of product formulations.

Overall, the successful operationalization of these ML models demonstrates the transformative potential of AI in the manufacturing industry. By bridging the gap between research and practical application, these models have the power to revolutionize processes, drive innovation, and create sustainable competitive advantages.

5. Hybrid modeling: A faster path to model maintenance

Calibrating TSWG models is often challenging because of the large number of kinetic and rate mechanism parameters involved, which typically exceed 25. This high dimensionality makes it difficult to attribute changes in model behavior to individual parameters under specific operating conditions. In this work, attempts at global calibration – that is, simultaneously fitting all parameters across all experimental runs – led to numerical instability and early termination, highlighting both identifiability and robustness limitations in the full-parameter formulation.

As a result, common practice is to perform reduced calibration by applying a sensitivity index threshold to select a subset of parameters and operating variables for estimation. In many cases, the liquid-to-solid (L/S) ratio is prioritized because of its strong correlation with the PSD. While practical, this strategy can compromise predictive accuracy and limit the model's ability to generalize across the operating envelope. Moreover, it is insufficient when other operating variables materially affect the dynamics of interest. For example, powder feed rate and screw speed are essential for capturing energy input and process intensity in twin-screw granulation and, therefore, must be included in calibration and subsequent prediction of CQAs.

To address these limitations, this chapter introduces a hybrid inverse-modeling strategy in which ML is first used to infer difficult-to-estimate (“rare”) kinetic parameters rapidly, after which the inferred parameters are coupled with the PBM to predict CQAs. The intent is to preserve the mechanistic structure while improving calibration speed, numerical robustness, and re-parameterization capability when operating conditions change.

The workflow (see **Figure 11**) begins by identifying the kinetic parameters that most strongly influence the target CQA. This step combines mechanistic insight (e.g., dominant aggregation, breakage, and layering pathways) with sensitivity analysis to prioritize parameters that materially affect the average particle size produced by the twin-screw granulator. Next, input–output datasets are generated for the selected parameters and operating variables. Given the complexity and cost of exhaustive experimentation, the present study did not rely on new experimental campaigns for data generation. Instead, datasets were assembled by leveraging published sources and/or industrial databases, supplemented by virtual experiments

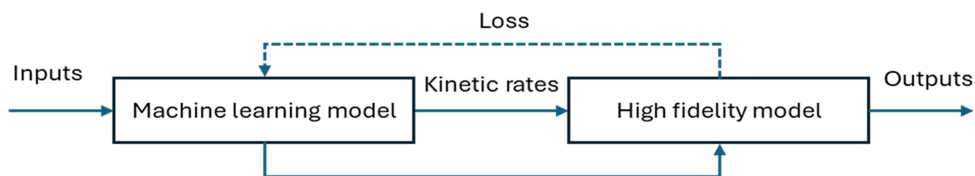


Figure 11.
Hybrid model (ML and PBM) algorithm structure.

generated in simulation software under controlled conditions to ensure systematic coverage of the calibration domain.

In the final step, the model’s performance was evaluated by simulating its predictions on the test data and comparing the results to the expected outcomes. Additionally, the simulation time was measured.

The ML model developed to predict breakage rate constants in the CE and KE, as well as the layering growth rate, exhibited a high degree of error, as shown in **Tables 8, 9** and **10**. The calculated breakage fraction of the fine kinetic parameter produced a negative result, indicating a physically unrealistic outcome. This anomaly arises from the model’s formulation, which does not incorporate constraints to enforce non-negativity. The high error can also be attributed to the inherent challenges of solving inverse problems. Unlike forward problems with well-defined, unique solutions, inverse problems may have high sensitivity to input variations. Small changes in the input data can lead to significant changes in the predicted parameters, making the model unstable. Our analysis also identified the layering maximum growth rate, the breakage rate constant in CE, and the breakage rate constant in KE as the most sensitive kinetic parameters. This sensitivity contributes to the high prediction errors observed in the model.

Despite these limitations, the model can still provide an initial estimate of the kinetic parameters, serving as a starting point for further refinement. Performance improvements for each model can be achieved by minimizing interactive effects between parameters, expanding the training dataset, and optimizing the model’s learning rate.

Surrogate model output			
Y1	Breakage rate constant on CE	39.4669	1.5
Y2	Breakage rate constant on KE	19.1587	2.5
Y3	Fraction of fine on CE	-0.6226	0.4

Table 8.
Twin-screw wet granulation breakage parameter surrogate model prediction vs. high-fidelity model calibration.

Surrogate model output			
Y4	Layering growth rate	2557.346	2666.5

Table 9.
Simulation stats: twin-screw wet granulation layering surrogate model prediction vs. high-fidelity model calibration.

Surrogate model output			
Y5	Pore penetration	0.68	0.52
Y6	Standard deviation of initial droplet size	0.712	0.40

Table 10.
Twin-screw wet granulation nucleation surrogate model prediction vs. high-fidelity model calibration.

Model output	Surrogate	High-fidelity model
Number of equations	10	54,271
Simulation time (s)	0.40672	83.7564

Table 11.
Simulation stats: ML surrogate model prediction and high-fidelity model calibration.

While PBMs are the traditional approach for modeling twin-screw granulation, this study explores the potential of ANNs as a complementary tool. PBMs excel at tracking particle property changes throughout the process, predicting key outputs like PSD, liquid content, and porosity. The developed ANN, however, demonstrates promise for rapid and reliable prediction (see **Table 11**), even if its accuracy is currently lower than that of established PBM methods.

This advantage in speed and reliability makes ANNs suitable for implementing MPC strategies. Furthermore, the ANN can be coupled with PBM models to accelerate the calibration process during continuous pharmaceutical manufacturing development.

5.1 Integrated workflow: AFRAME

The standard workflow focuses primarily on the R&D phase of model building. However, for the model to be truly valuable, it needs to be deployed online for real-world use. Once the designed experiments have been performed and the chosen measurements have been collected and imported, they can be used to generate better estimates of the parameters. The accuracy of the model and the quality of the estimated parameters can be quantified through cross-validation and uncertainty propagation analysis. If necessary, the model with the updated parameters can then be used to design further experiments, and the cycle is repeated until the desired level of accuracy is attained. The standard workflow for model development and parameter estimation has some limitations. Iterative experimentation and analysis can be a slow process, especially when dealing with complex models and large datasets. The selection of experimental designs and the interpretation of results can be influenced by human bias. Limited exploration: Traditional methods might not explore the full design space efficiently, potentially missing optimal solutions.

While the conventional workflow is often adequate for research and development, real value is realized only when models are deployed online for operational decision-making (e.g., real-time prediction, monitoring, optimization, or control). Online deployment introduces additional requirements that must be addressed explicitly, including numerical robustness and sensitivity to data drift as materials and equipment evolve. Accordingly, adaptive framework for robust and accelerated model evaluation (AFRAME) is designed not only to streamline model development in the R&D phase but also to support sustainable online use by facilitating rapid re-calibration, structured updating with new data, and more reliable model behavior across the defined operating envelope.

The provided workflow (see **Figure 12**) illustrates an accelerated framework designed to bridge the gap between complex physical simulations and real-time digital operations through a tiered modeling strategy. At its core, the normal cycle establishes a reliable physical baseline by transforming raw sensor and experimental data into a validated mechanistic model via parameter estimation. This foundation is

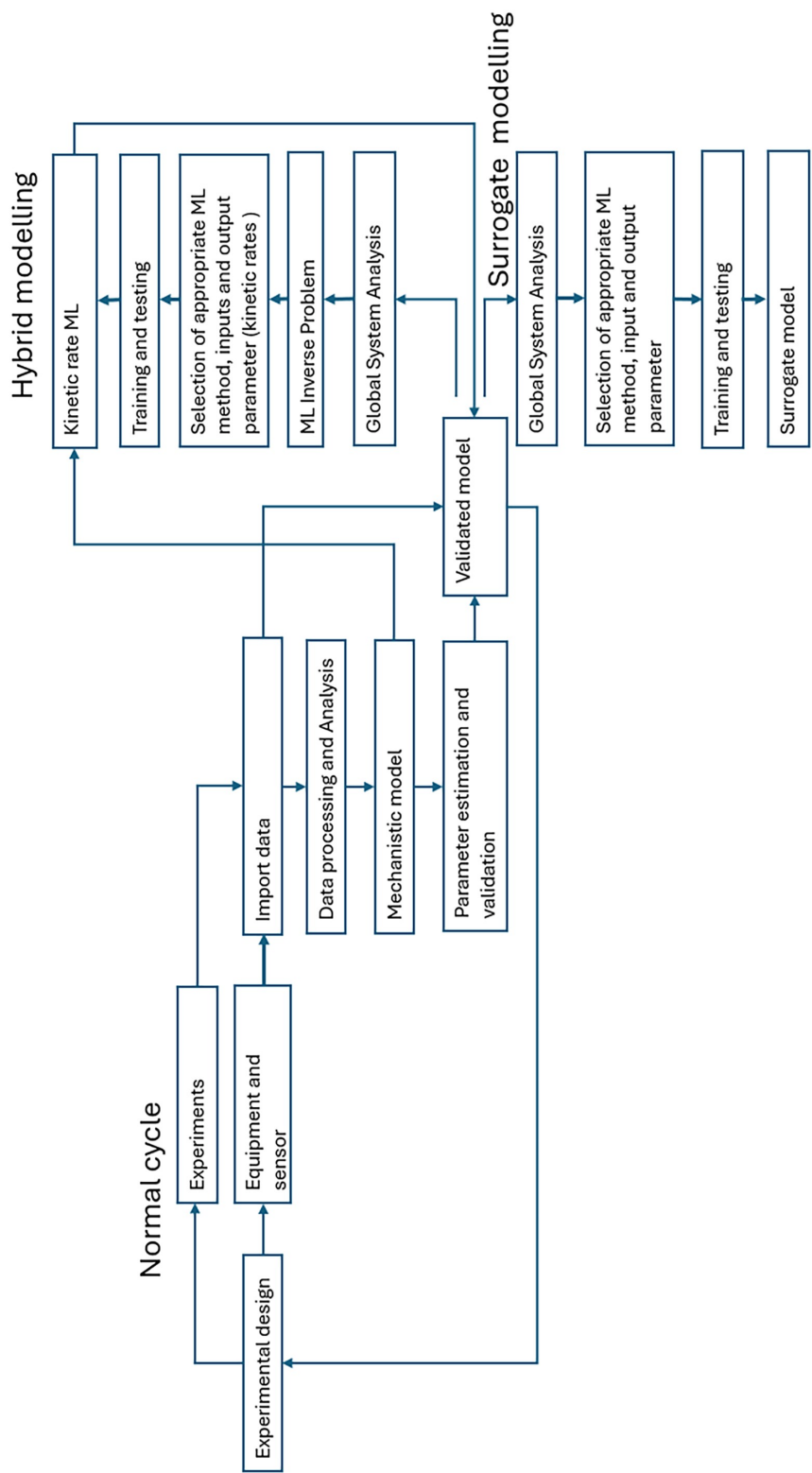


Figure 12. Proposed accelerated framework for the deployment of advanced models for digital operation (AFRAME).

augmented by hybrid modeling, which employs ML to resolve high-dimensional or poorly understood parameters – such as kinetic rates – by solving ML inverse problems and utilizing global system analysis to ensure the model remains both physically grounded and empirically accurate. The workflow culminates in surrogate modeling, which converts these computationally intensive models into lightweight, high-speed ML emulators, enabling the low-latency response required for DT applications. Crucially, the feedback loops from both the validated and surrogate models back to experimental design facilitate automated model maintenance, creating a self-correcting system that identifies data gaps and adapts to operational shifts without manual intervention.

6. Conclusion(s)

Developing an accurate process model remains one of the most time-consuming and resource-intensive steps in building and sustaining a DT. Although rigorous first-principles models can now be formulated for many unit operations, their direct use in digital operations is often constrained by two practical limitations: (i) computational cost that precludes repeated simulation, optimization, or real-time deployment; and (ii) limited numerical robustness across the full operating envelope. These constraints are particularly acute in online applications such as optimization and model predictive control, where predictions must be both fast and reliable. Hybrid modeling offers a pragmatic route to address these barriers by combining physics-based structure with data-driven approximations. In this chapter, a PBM framework was integrated with ML to preserve mechanistic interpretability while improving deployability.

A calibrated mechanistic model can be used as a high-fidelity “data generator” to create pseudo input–output datasets through systematic sampling. These datasets enable the training of surrogate models that execute rapidly and robustly within a defined domain while retaining the predictive fidelity of the underlying physics-based simulator. Using this strategy, the chapter developed TSWG surrogates derived from the calibrated model. The resulting surrogate comprised only 10 equations, achieved an R^2 of 0.998 for PSD prediction, and required 0.41 s per simulation. On independent test data, the surrogate reproduced key CQAs with simulation times below 5 seconds, supporting its suitability for DT integration and control-oriented studies.

The chapter also presented a hybrid PBM–ANN architecture in which ML provides rapid first-pass estimates of kinetic rate parameters within the calibration domain, which are then passed to the PBM to generate PSD predictions. This approach can accelerate calibration and enable rapid re-parameterization as operating conditions change. However, the results also underline important limitations: inverse estimation of kinetic rates – particularly for breakage and layering – may be ill-posed, nonunique, and sensitive to small perturbations in inputs. Performance improved as training coverage increased; for example, expanding the layering training set from 100 to 250 samples materially increased predictive accuracy, as reflected by the reported gains in R^2 .

Maintaining DT relevance over the process lifecycle requires continual model updating for two reasons: incorporating new data to improve predictive performance and adapting to process drift as equipment and materials evolve. A balanced strategy – combining automated updating with appropriate expert oversight – is therefore essential. Hybrid modeling can reduce time-to-deployment, especially when complete first-principles knowledge is unavailable, while reusable mechanistic structures (such as PBMs) support transfer to related processes with modest reconfiguration of ML components. Nevertheless, the addition of new process variables (e.g., via additional sensors) and associated retraining should be justified through a cost–benefit assessment that weighs expected improvements in predictive accuracy against the burden of data collection, model redevelopment, validation, and computational resources.

Conflict of Interest

The authors declare no conflict of interest.

Nomenclature

Symbols	Description	Units
a, a', b, b', c, c'	Conveying and kneading element breakage kernel empirical parameters	-
B_{M1} and B_{M2}	Breakage functions	-
DYS	Dynamic yield stress	kPa
d_{50}	Median particle diameter	μm
$G_i(v)$	Layering growth rate	$\mu m.s^{-1}$
$G_i(x)$	Layering linear growth rate	$\mu m.s^{-1}$
$k_{i, comp}$	Consolidation rate constant in compartment i	s^{-1}
$\dot{m}_g$	Granule mass flow rate	$kg.h^{-1}$
$\dot{m}_p$	Powder mass flow rate	$kg.h^{-1}$
M_i	Total mass of fine powder, granules, and liquid in the compartment i	kg
$P_{b, i}(v)$	Breakage fragment distribution function volume v in compartment i	-
PFN	Powder feed number	-
$R_i(\epsilon)$	Consolidation rate in compartment i	s^{-1}
s^*	Droplet pore penetration	-
V_{drop}	Single drop volume	m^3
$x_{i, v}^{in}$ and $x_{i, v}^{out}$	Mass fraction of vapor phase in the inlet and outlet stream	$kg.kg^{-1}$
x_{lc}	Lower bound granule diameter	mm
x_{uc}	Upper bound granule diameter	mm
$y_{w, i}$	Moisture content in compartment i	%
β	Weibull function scale parameter	-
γ	Weibull function shape parameter	-
μ_{drop}	Drop size	mm
$\mu_{drop, mean}$	Mean droplet size	mm
σ_{drop}	Drop size standard deviation	mm
$\sigma_{i, T, y}^2$	Total variance attributed to factor I on response y through all factors	-
$\sigma_{i, y}^2$	Variance attributed to factor i on response y	-
τ_i	Dispersion time constant in compartment i	s
ϕ	Granule mass fraction	-

Author details

Mohammad Zandi* and Donald Ntamo
University of Sheffield, Sheffield, UK

*Address all correspondence to: m.zandi@sheffield.ac.uk

IntechOpen

© 2026 The Author(s). Licensee IntechOpen. This chapter is distributed under the terms of the Creative Commons Attribution License (<http://creativecommons.org/licenses/by/4.0/>), which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited. 

References

- [1] Lee SL, O'Connor TF, Yang X, Cruz CN, Chatterjee S, Madurawe RD, et al. Modernizing pharmaceutical manufacturing: From batch to continuous production. *Journal of Pharmaceutical Innovation*. 2015;**10**(3):191–199. DOI: 10.1007/s12247-015-9215-8
- [2] Seem TC, Rowson NA, Ingram A, et al. Twin screw granulation — A literature review. *Powder Technology*. 2015;**276**:89–102
- [3] Teżyk M, Milanowski B, Ernst A, Lulek J. Recent progress in continuous and semi-continuous processing of solid oral dosage forms: A review. *Drug Development and Industrial Pharmacy*. 2016;**42**(8):1195–1214. DOI: 10.3109/03639045.2015.1122607
- [4] Javaid M, Haleem A, Singh RP, Suman R. An integrated outlook of Cyber-Physical Systems for Industry 4.0: Topical practices, architecture, and applications. *Green Technologies and Sustainability*. 2023;**1**(1):100001. DOI: 10.1016/j.grets.2022.100001
- [5] Barrasso D, El Hagrasy A, Litster JD, Ramachandran R. Multi-dimensional population balance model development and validation for a twin screw granulation process. *Powder Technology*. 2015;**270**(PB):612–621. DOI: 10.1016/j.powtec.2014.06.035
- [6] Cameron IT, Wang FY, Immanuel CD, Štěpánek F. Process systems modelling and applications in granulation: A review. *Chemical Engineering Science*. 2005;**60**(14):3723–3750. DOI: 10.1016/j.ces.2005.02.004
- [7] Iveson SM, Litster JD, Ennis BJ. Fundamental studies of granule consolidation part 1: Effects of binder content and binder viscosity. *Powder Technology*. 1996;**88**(1):15–20. DOI: 10.1016/0032-5910(96)03096-3
- [8] Wang LG, Morrissey JP, Barrasso D, Slade D, Clifford S, Reynolds G, et al. Model driven design for twin screw granulation using mechanistic-based population balance model. *International Journal of Pharmaceutics*. 2021;**607**:120939
- [9] Litster J, Bogle IDL. Smart process manufacturing for formulated products. *Engineering*. 2019;**5**(6):1003–1009. DOI: 10.1016/j.eng.2019.02.014
- [10] Rogers AJ (2015) Process systems engineering methods for the development of continuous pharmaceutical manufacturing processes. PhD thesis, Rutgers, The State University of New Jersey (RUcore record: rutgers-lib/47570)
- [11] Metta N, Ghijs M, Schäfer E, Kumar A, Cappuyns P, Van Assche I, et al. Dynamic flowsheet model development and sensitivity analysis of a continuous pharmaceutical tablet manufacturing process using the wet granulation route. *Processes*. 2019;**7**(4):234. DOI: 10.3390/pr7040234
- [12] Ekins S. The next era: Deep learning in pharmaceutical research. *Pharmaceutical Research*. 2016;**33**(11):2594–2603. DOI: 10.1007/s11095-016-2029-7
- [13] Siemens Digital Industries Software. gPROMS FormulatedProducts Documentation (Add Software Version/build and

Document ID if Available). 2024
Version 2023.2.1

[14] Ntamo D, Montero E, Omar C, Highett M, Moss D, Soulam O, et al. Industry 4.0 in action: Digitalisation of a continuous process manufacturing pilot plant. *Digital Chemical Engineering*. 2022;**3**:100025

[15] Davidopoulou N, Ouranidis A, Artemaki S, Vavilis T, Bikiaris DN. Digital twins and artificial intelligence in Pharma 4.0: A review. *Pharmaceutics*. 2022;**14**(10):2113. DOI: 10.3390/pharmaceutics14102113

[16] Matsunami K, Miura H, Yaginuma Y, Tanabe H, Badr S. Surrogate model-based dissolution profile estimation based on dissolution mechanisms. *Computers and Chemical Engineering*. 2023;**176**:108141. DOI: 10.1016/j.compchemeng.2023.108141

[17] Rogers AJ, Hashemi A, Ierapetritou MG. Modeling of Particulate Processes for the Continuous Manufacture of Solid-Based Pharmaceutical Dosage Forms. *Processes*. 2013;**2**:67–127. DOI: 10.3390/pr1020067

[18] Vusiness Wire *PSE: Sanofi joins Systems-based Pharmaceuticals Alliance, 2019. Press Releas*. Available from: <https://www.sttinfo.fi/tiedote/69871239/pse-sanofi-joins-systems-based-pharmaceutics-alliance?publisherId=58763726> [Accessed: 2026-March]

[19] Saltelli A, Annoni P, Azzini I, Campolongo F, Ratto M, Tarantola S. Variance based sensitivity analysis of model output. Design and estimator for the total sensitivity index. *Computer Physics Communications*. 2010;**181**(2):259–270. DOI: 10.1016/j.cpc.2009.09.018

[20] International Society for Pharmaceutical Engineering (ISPE). GMP resources. ISPE Regulatory Resources (GMP). 2024

[21] Wang LG, et al. (2020) ‘A breakage kernel for use in population balance modelling of twin screw wet granulation’, *Powder Technology*, **363**, pp. 525–540. DOI: 10.1016/j.powtec.2020.01.024

[22] Klein P, Bergmann, R. (2018). Data Generation with a Physical Model to Support Machine Learning Research for Predictive Maintenance. *Lernen, Wissen, Daten, Analysen*

[23] A Saxena, K Goebel, D Simon, N Eklund, Damage propagation modeling for aircraft engine run-to-failure simulation, in: 2008 international conference on prognostics and health management, IEEE, 2008, pp. 1–9

[24] Wang LG, Omar C, Litster J, Slade D, Li J, Salman A, et al. Model driven design for integrated twin screw granulator and fluid bed dryer via flowsheet modelling. *International Journal of Pharmaceutics*. 2022;**628**:122186