

EVOLUTIONARY OPTIMIZATION OF BATCH PROCESS SYSTEMS USING IMPERFECT MODELS

E C MARTINEZ* AND J A WILSON†

**Instituto de Desarrollo y Diseño, Consejo Nacional de Investigaciones Científicas y Técnicas, Avellaneda 3657 - Santa Fe, S3002 GJC (Argentina)*

†School of Chemical, Environmental and Mining Engineering, University of Nottingham, University Park, Nottingham NG7 2RD (England)

(Received 27 June 2002; Accepted 04 September 2002)

Finding optimal operating conditions fast for batch process systems is a key competitive factor in fine chemicals and pharmaceutical industrial sectors to increase the value added in many stages of a product/process lifecycle, including recipe development, scale-up and run-to-run optimization. However, specialty chemical processing increasingly features products with short market windows that make the development of a detailed model unattractive in terms of both time and economy. Moreover, key measurements of process performance are sparse, noisy and delayed. As a result, operating procedures for batch processes are often developed rather conservatively and scale-up is often carried out with a significant loss of profitability. In this work, a comprehensive review of the concepts, techniques and issues involved in the simultaneous development and use of imperfect models for batch process optimization is provided. The role of active learning for model-based evolutionary optimization in the face of modelling errors is highlighted. To deal with scarce and biased data compounded with batch-to-batch variations a statistical characterization of optimum operating conditions is discussed. The problem of model discrimination is addressed using nonparametric statistical methods and pair-wise bisection ranking of data points. The Kendall's tau statistics is used to measure model *correlation* and *concordance* with regards to alternative optimum predictions. Less promising model hypotheses are gradually eliminated using and exponential "cooling" procedure that trades off *exploitation* with *exploration*. A model-based evolutionary optimization methodology that accounts for process-model(s) mismatch and end-point constraints is proposed. Potential applications to help planning experiments during process development and scale up stages are discussed.

Key Words: Process Development; Batch Processes; Experimental Design; Modelling for Optimization; Non-Parametric Statistics; Evolutionary Optimization; Learning; Model Discrimination

1 Introduction

Finding optimal operating conditions fast for batch process systems is a key competitive factor in fine chemicals and pharmaceutical industrial sectors to increase the value added in many stages of a product/process lifecycle, including recipe development, scale-up and run-to-run optimization. Batch processes are associated with a significant portion of the value delivered by chemical manufacturing industries. To deploy and sustain the required value chain in this competitive environment, modelling methodologies that allow faster and less costly process development and scale-up of an ever-increasing number of products are needed. Product customization and reduced time-to-market are features of central concern to speed up

the process development lifecycle in order to succeed against an increasing number of alternative products^{1,2}. Also, the optimal operation of batch units is becoming of great importance, as stringent product quality limits and reproducibility requirements are standards that should be satisfied for product acceptance. However, industrial practice still relies heavily on past experience and human judgement, which often mean sub-optimal operating conditions and a considerable loss of profitability. These industrial needs have attracted a great deal of attention from academic researchers towards the development of methodologies for process modelling, advanced control and optimization of batch processes.

The use of a model has been a recurrent feature in most of the proposed methods. However, specialty chemical processing increasingly features products with short market windows that make the development

*Fax: +54 342 4553439; email: ecmarti@ceride.gov.ar

†Fax: +44 115951 4115; email: J.A.Wilson@nottingham.ac.uk

of a detailed model unattractive in terms of both time and economy. Moreover, key measurements of process performance are sparse, noisy and delayed. There exists an extensive literature on model-based optimization techniques for batch processes. There are two idealized assumptions often made with regards to the availability of the information and knowledge required for model development and use. One of the idealized situations is the perfect model condition³. In this situation, it is assumed that it is possible to resort to a perfectly accurate and sufficiently detailed model comprising all thermodynamic and kinetic relationships required to describe the batch process behaviour. Under this postulate there is no need to measure and feedback any measurement from the batch process. The optimal operating policy derived from the "perfect" model is applied "open-loop" since it is robust enough to compensate for any source of process variability and anticipate or measure the effect of any disturbance. The second idealized situation is the comprehensive instrumentation condition. This condition refers to the case where all state variables of the process can be measured on-line with sufficient accuracy and frequency. If everything that matters can be measured, it can be easily argued that the need for a first principles model is lessened. If we can measure everything, an entirely inductive data-driven model can be readily developed directly from data. Unfortunately, neither of these assumptions is valid in batch industrial environments⁴⁻⁶. Thus, migration from laboratory conditions to production runs is often made with high levels of uncertainty and risk. Modelling for optimization^{6, 7} has become of crucial concern as the limited number of costly runs performed during process development and scale-up leave little room for the development of a detailed model of the process.

One central concern of modelling for optimization is how experimental design in the modelling lifecycle can be best approached considering poor knowledge about phenomena involved, poor reproducibility of run outcomes and modelling bias inherited from lab scale experimentation. Modelling for optimization constitutes a shift from the traditional approach which assumes a separation between a model and its use, to simultaneous model development and process optimization. In the latter case, the purpose of modelling is focusing experimentation to help

defining an optimum operating policy. Different model hypotheses comprising structure and parameterization upon data are postulated, refined and discharged in the path to the process optimum. Thus, each new experiment in the sequence of runs must be able to bring new information to increasingly reduce the uncertainty regarding optimum location.

In modelling for optimization, it is convenient to refer to two building blocks for model identification: (i) *first principles* hypotheses about the inner workings (conservation equations, constitutive laws, etc.) of the process and, (ii) *data* that characterize the actual (observed) behaviour of the process under study through parameter estimation. As a result, the modelling process for the batch process cannot be entirely knowledge-driven or data-driven alone. Instead, the strategy for model development should necessarily be one that takes the best of both worlds. One avenue for doing this in a reaction system (the most ubiquitous for batch systems), initially proposed by Filipi *et al*^{8, 9}, is *Tendency Modelling*. A "tendency model" is a low order, nonlinear, dynamic model that approximates the stoichiometric and kinetic relationships of a process using the available plant data along with fundamental knowledge of the process characteristics. The model structure and parameters are incrementally updated as more data becomes available. The main use of a tendency model is to determine a direction towards the optimum. For this to be feasible, the model should be able to extrapolate to operating conditions quite different from those used for model identification. First principles and constitutive laws incorporated whilst building the model hopefully provide this capability by constraining the degrees of freedom available for data fitting.

Since Tendency modelling was first proposed^{8,9}, a number of new developments and successful applications have been reported (e.g. Uhlemann *et al.*¹⁰). The issue of model structure identification from process data has attracted a great deal of attention¹¹⁻¹⁸. Singular value decomposition^{11,12} and Structured Target Factor Analysis^{14,16,19,20} have proved to be very successful tools for many examples. Also, a methodology to quantify the impact of model parametric uncertainty onto the uncertainty of the predicted optimal operating condition has been developed¹⁶. However, the issues of experimental design for optimization and model discrimination and selection are still open problems²¹.

In order to achieve the goal of optimal operation of batch processes in the face of uncertainty, a number of requirements are imposed on modelling for optimization for it to be successful. The first is how data gathering is increasingly biased towards more profitable operation bearing in mind that model adaptation (parameters and structure) is simultaneously occurring with evolutionary optimization. Also, the characterization of optimal operation using process data that are both scarce and noisy must be addressed. Another important issue is effective discrimination among competing model structures considering that we are not interested in the model *per se* but in a more profitable operation. Successfully addressing these issues is mandatory to define a sequential experimental design strategy that can systematically reduce the uncertainty about the location of the process optimum.

This contribution is organized as follows. To emphasize the role of modelling as a tool for process optimization in the scope of industrial applications, Section 2 discusses the main characteristics of batch processes. Later on, Section 3 focuses in the modelling for optimization cycle. To deal with scarce and biased data compounded with batch-to-batch variations a statistical characterization of optimum operating conditions is presented in Section 4. The problem of model discrimination is addressed using nonparametric statistical methods and pair-wise bisection ranking of data points. The importance of active learning to compensate for modelling errors is discussed in Section 5. Finally, Section 6 proposes a model-based evolutionary optimization methodology that accounts for process-model(s) mismatch and end-point constraints.

2 Characteristics of Batch Processes

The industrial setting of batch processes exhibits a number of characteristics with regard to objectives, dynamic behaviour, scarce data and noisy measurements that constitute a challenging problem for model-based optimization of batch processes^{4, 6, 22, 23}.

2.1 Objectives

Industry sees the potential of process modelling primarily with respect to continuous improvement, guaranteeing product value and reproducibility, speeding up scale up and safeguarding. Increasing productivity of a processing unit is also very desirable since reducing the duration of a batch allows lower installed capacity or higher plant throughput.

2.1.1 Reproducible Products

Batch processes, including pharmaceutical and biochemical production, commonly involve significant variability in performance. Continuously changing levels of impurities in raw materials and poor understanding of process behaviour are often targeted as the root causes of variability. In the first case, material recycles give rise to batch-to-batch variations in the initial state of each run. On the other hand, operating procedures and control loops that are not designed upon a proper representation of the phenomena involved introduce intra-run variations that have a negative effect in reproducing product attributes.

Variability is of particular significance for processes involving polymerizations or crystallizations where there is no possibility of downstream blending to achieve the desired product specifications. Reproducibility of product attributes is even more critical for pharmaceuticals, where decisions about the optimal way of operating the process should be taken early on during process development at the laboratory or pilot-plant scale. Because of *FDA* regulations, optimization of operating conditions is strictly prohibited during production.

2.1.2 End-use Product Properties

The value added in specialty and pharmaceutical sectors is set by specific properties of the final product that are related to its intended use²⁴. These properties are directly or indirectly related to the molecular properties of the product. Examples of these properties include biodegradability, solubility, tensile strength, melt index, scrub resistance, film adhesion, oil resistance, etc. Typically, each end-use property is required to lie in a target range, which is mandatory for product acceptance. For example in the “Setal” process²⁵, end-use properties of the resulting alkyd resin are acid number, viscosity and hydroxyl number of the product.

2.1.3 Safety

Of paramount concern when one migrates from the bench or laboratory scale to the industrial size units is the possibility of unsafe operating conditions, such as runaway situation or venting health threatening substances into the air. This is caused by our imperfect knowledge of the inner working of the process. Typically, the production of specialty and fine chemicals involves complex reaction systems where many additional and unknown reactions

might occur, depending on the operating conditions. On the other hand, environmental considerations have tightened the level of uncertainty and risk that are considered tolerable on the shop floor when new processes are put into operation.

2.1.4 Scale-up

The global competitive climate is imposing on most products shorter lifecycles and, to take advantage of innovation, shorter time-to-market. Thus, product changeovers are frequent and scale-ups from laboratory data directly to production plants are seen as big opportunities to avoid costly pilot-plant studies. Again, uncertainty and risk are of central concern bearing in mind that process behaviour is poorly understood.

2.1.5 Productivity

Multipurpose batch reactors are expensive setups operated by qualified personnel. Operating conditions that can sufficiently reduce the duration of a batch run so as to make room for processing more runs per day or week are productivity levers for the complete production line. Also, expensive reactants charged should be converted as much as possible into the desired product while minimizing loss to undesired side products which often degrade product quality and demand further separation costs.

2.2 Behaviour

Dynamic behaviour of batch processes features certain characteristics that make them more difficult to model and operate in comparison with continuous processes. To a large extent the lack of automation in many industrial environments^{4,6} can be explained on the basis of poor modelling of batch process behaviour comprising an evolution from an initial, partially known state to a very different final state. This transformation changes the chemical composition of the system significantly and with that the rate of evolution, which exhibit very nonlinear dependencies on species concentrations and temperature. As a result, all the associated physical parameters, such as viscosity, vapour pressure, specific heat content, etc., are dramatically changed while the process evolves. Thus, process behaviour is inherently non-linear and time-varying due to disturbances entering into the process either in a batch-to-batch form or while a run is being processed.

Batch process behaviour often unfolds in an irreversible way thus limiting the window of opportunity for taking corrective actions. The ability to influence the attributes of the endpoint state typically decreases as the run progresses. For history-dependent end-use properties such as shape and size distributions of crystals or polydispersity in polymers, if a corrective action is not taken very early on in the run there is little room to introduce later remedial corrections that may bring the final product back to specifications.

2.3 Measurements

Despite recent developments in chromatographic and spectroscopic sensors, on-line measurements of chemical compositions are still rare in industrial reactors. Commonly, chemical composition of the reactor content is determined by drawing samples that are analyzed off-line with considerable delay. Typically, sample results are known only after the run is finished, which leaves little room for taking corrective actions. Readily available measurements are commonly physical quantities such as temperature, pressure, pH, ORP, turbidity, electric conductivity or torque. These measurements are much less informative about the evolution of the batch and exhibit low reliability due to a wide range of different conditions that need to be covered. As more advanced sensors for inferring chemical compositions and reactions rates are becoming less expensive and more reliable, the limitation of appropriate measurements may be lessened in the years to come²⁶.

3 Modelling for Optimization

W. Edward Deming, the father of modern quality control, is reputed to have made once the following remark: "*All models are wrong; some models are useful for a purpose.*" Deming stresses the importance of *model valuation* in terms of utility (profitability, risk, etc.) to the owners of the model. Valuation is also related to the cost of the modelling process, i.e. measurements and experiments. Modelling for optimization focuses on model valuation under uncertainty by emphasizing that process models are means rather than ends in themselves.

3.1 Problem Statement

A process model is understood here as a mathematical device comprising of three basic hypotheses: *process invariants*, *constitutive "laws"*

and *fitting parameters*. Process invariants and constitutive laws make up the structure of the model. Even though some type of prior knowledge can be used for stating each model hypothesis, for batch processes they are mostly derived from experimental data. Process invariants represent structural constraints for process behaviour such as the number of independent reactions in a reactor^{18, 27-29}. These invariants correspond to the enforcing of conservation equations and topological constraints that confine process state evolution to a given sub-space. Constitutive laws, in turn, are functional relationships between variables, e.g. the power law for reaction rates, *gas* behaviour equations and *Raoult's Law* for vapour-liquid equilibrium. Typically, these relationships are only valid for a given operating regime, which constraint the extrapolation capability of a model. Finally, constitutive laws provide some room for a few model parameters that can be estimated by fitting to experimental data.

The combination of process invariants and constitutive laws with process data results in a *Grey-box* model that exhibits better predictive capabilities than a pure inductive or black-box model (e.g. a neural network) because it incorporates some fundamental process understanding through model structure¹⁵. Model structure acts as a *shaping* factor that diminishes the available degrees of freedom for model parameterization in a meaningful way. In a chemical reactor, for example, model structure is comprised of stoichiometric and kinetic considerations. The stoichiometric component is constructed by identifying a small set of independent overall reactions that adequately describe the compositional changes in the observed process data. Linear algebra techniques are used in the systematic derivation of a set of reactions or "factors" that account for most of the compositional variation in the data. This results in a set of reaction invariants, which along with rate laws constitute a model structure. It is worth noting that, even though model structure is based on first principles and constitutive laws, its identification is heavily dependent on process data.

Any useful model of a chemical process must be assessed with regards to a purpose. In 'modelling for optimization' the model is merely a *means* to a clear-cut *end*, namely to help systematically improve and optimize process operating conditions despite poor understanding of process behaviour

and uncontrollable disturbances entering into the process. The process model here is not an end in itself as it is in kinetic studies³⁰. The main objective of building a model is instead to help locate the optimum operating condition as precisely as possible with minimum experimental effort. The strategy chosen to obtain sequentially the experimental data needed for this purpose will influence greatly the number of modelling runs, the cost involved and the length of time required to accomplish the objective stated above. Ideally, operating conditions need to be progressively biased towards the most profitable operating conditions for the process. However, due to modelling errors, optimal operation can only be achieved if exploration of alternative model hypothesis is also systematically practiced. Therefore, the generation of data from the batch process should be planned so as to bring the most meaningful information from each experimental run. It is worth noting that as the model gradually evolves from one based on first principles and constitutive laws to a more inductive (data-driven) model, modelling for optimization 'gracefully degrades' to a more costly trial-and-error approach.

As shown in Fig. 1, modelling for optimization proposes an entirely different approach for relating process modelling with model-based optimization. Traditionally, a model is developed first and, only once the model has been validated, its optimization undertaken. Consequently, a great deal of experimentation is done to obtain data all over the region of interest that permit one to derive and test model hypotheses. This course of action is acceptable only if the model structure is good enough and data gathering is not too costly in terms of time and money. However, due to an incomplete understanding of the inner workings of the process and limited room for exploring risky or low performance operating conditions, model development should be tightly integrated within an inner loop that also includes active learning as shown in Fig. 1. Active learning seeks to influence the generation of data by establishing a two-way relationship between the process and the modeller (see Fig. 2). The learning block allows changing model hypotheses *a posteriori* to account for lack of improvement or even wasted batches. Model discrimination and the trade off between exploration and exploitation in model selection are central components of modelling for optimization. Also,

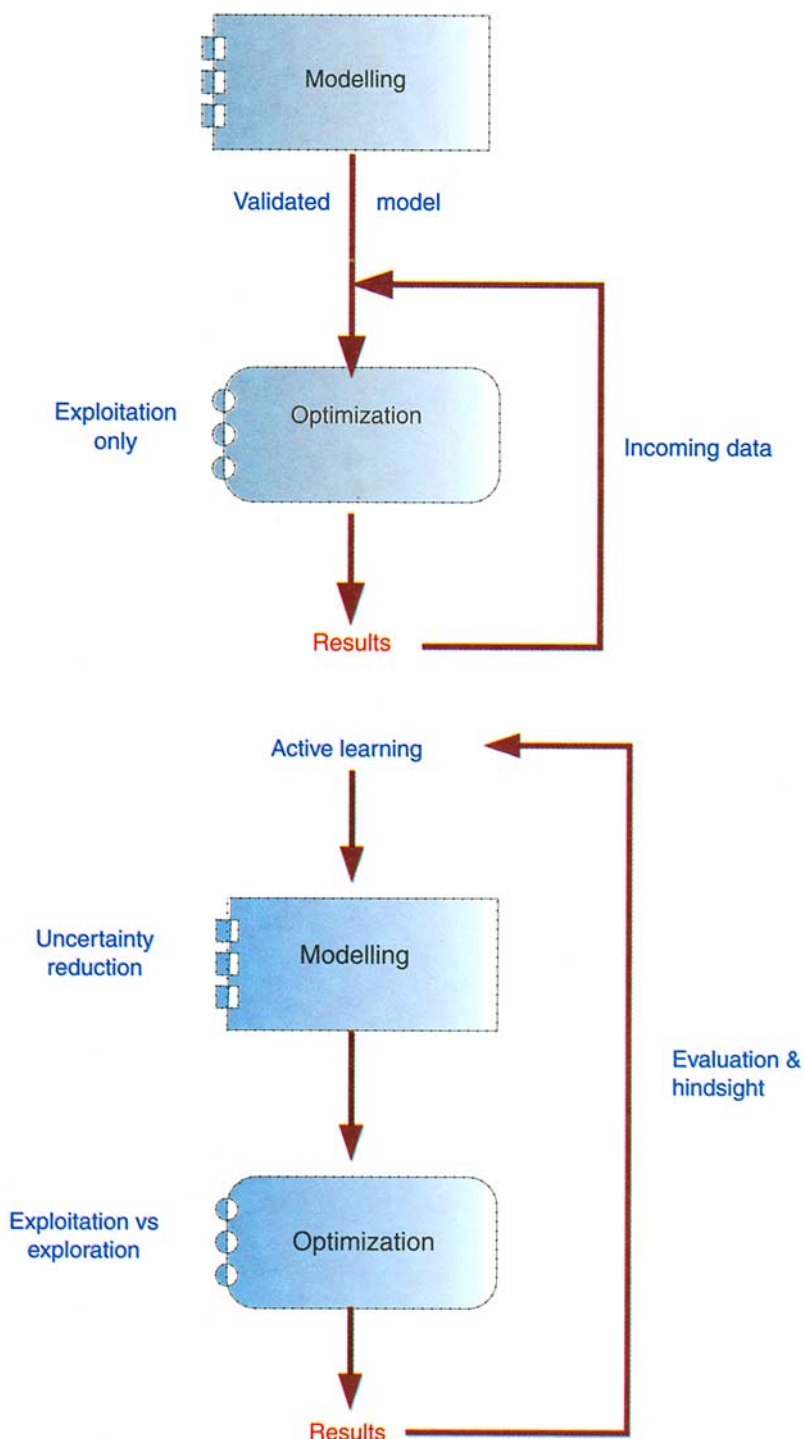


Fig. 1(b) Modelling-for-optimization approach

the issue of dealing with noisy and biased data should be addressed.

3.2 Model Identification from Process Data

Fig. 3 summarizes the various steps for identifying candidate process models from an initial set of data. First, the number of reactions to explain

the observed composition dynamics^{11, 12, 18, 27, 28} is determined. For this purpose, singular value decomposition and target factor analysis are used to identify a number of uncorrelated factors that explain the variance (composition changes) in the data set¹¹⁻¹⁸. The decomposition allows for easy identification of reaction invariants describing the

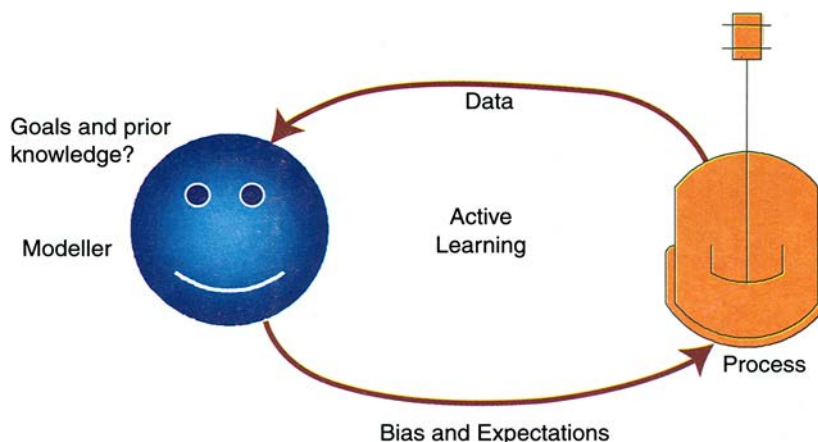


Fig. 2 Active learning from process data in the model development lifecycle

stoichiometric space of the observed data. The extents of reactions are then tested to identify plausible stoichiometric models. For each stoichiometry, accompanying kinetic model(s) must be constructed. To this end, the functional form and parameters of the rate law for each reaction must be chosen. The rate law can be of several types including the power law kinetics, inhibition kinetics or more complex kinetic models such as the *Monod* kinetics. Model parameters can consist of activation energies, pre-exponential factors, rate orders, etc., which are estimated upon process data. Finally, each alternative model is assessed to pinpoint predicted behaviours that may invalidate it, e.g. an activation energy having a negative value.

The capability of process models to account for the qualitative behaviour of the process as well as the approximate quantitative behaviour is heavily dependent on the richness and locality of the operating conditions in the data set. This means that model(s) extrapolation capabilities are much limited by what information and knowledge can be actually derived from the available data set. For example, consider the case of a side reaction that takes place at higher temperatures than those included in the experimental data set. The only way competing models can account for the profitability loss resulting from this undesirable reaction is experimenting in operating conditions where the product is contaminated or reactant losses are significant enough. Also, locality of model predictive capability can make it fail to capture the effect of transport limitations or mixing characteristics. Hence, choosing the operating conditions for the next experiment gives rise to a crucial dilemma: exploitation versus exploration. To optimize the process it can only be resorted to the exploitation of the knowledge that is already

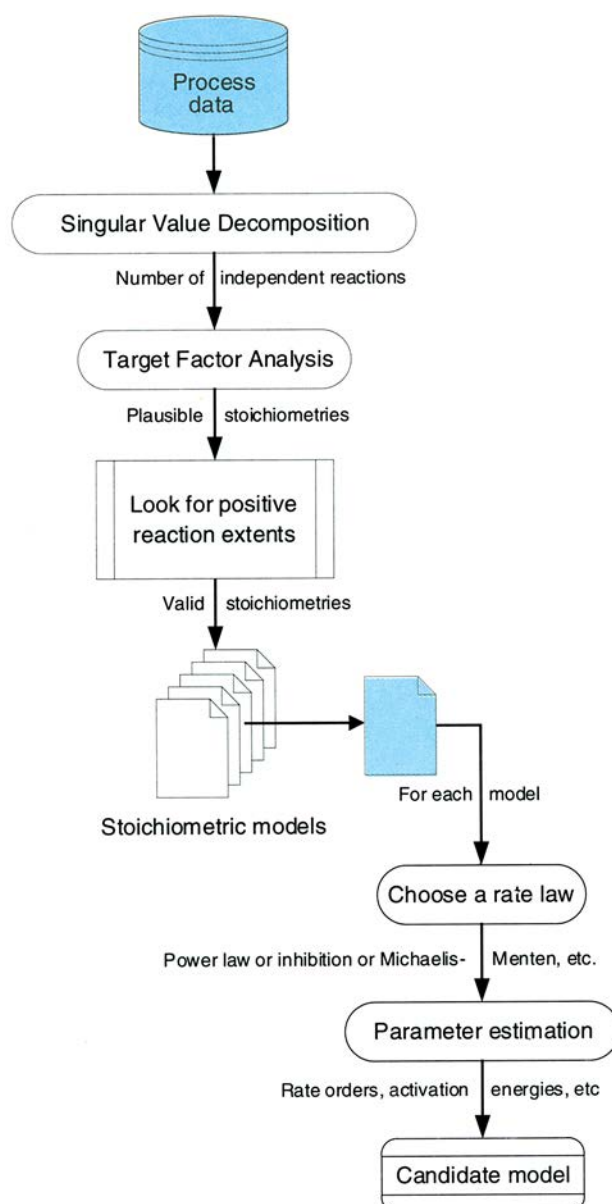


Fig. 3 Generating candidate models for reaction systems from process data

available in the current data set, whereas exploration seeks to enrich the data set so as to have more (new) knowledge to exploit further in the future⁷. This dilemma is intrinsic in resorting to imperfect models for process optimization and can be addressed by introducing a learning dimension in process modelling.

3.3 The Modelling for Optimization Cycle

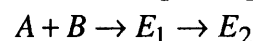
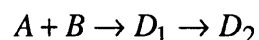
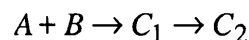
Each process model features a structure and set of parameters that are heavily influenced by quantity, quality and locality of data. Modelling for optimization is an experimental design strategy where alternative model hypothesis are postulated, tested and re-formulated in a sequence of experimental runs. As more data is gathered, poor improvement or even wasted batches call for changes and refinement of model hypotheses. Thus, model discrimination and selection are always key issues when deciding the operating conditions for the next experiment. As long as significant process/model mismatch exists there is room for exploring further, so the appropriate steps should be taken to actively seek the generation of additional process data that can help to discriminate among competing models and find better operating policies. Hence, *concordance* or *rivalry* among alternative model hypotheses with regards to their capability for locating process optimum is a key issue when dealing with a limited amount of experimental data.

The reward (profit improvement) obtained in each experimental run reinforces some future model selections and discourages others. Also, as new data points are incorporated into the data set, structure and parameters of alternative models will change as a result. Should the model with the highest (predicted) improvement always be followed, sampling bias will be introduced in the data set. It is advisable, mainly in initial stages of experimentation, to compensate for this by providing some room to follow less promising models whose optimal operating policy has a smaller predicted improvement. Due to noisy data and modelling errors chances are that one of these models may give rise to operating policies with an unexpectedly high reward. The modelling for optimization cycle is summarized in Fig. 4. From an initial data set, alternative model hypotheses are formulated. Following model discrimination and selection, the

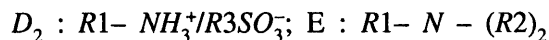
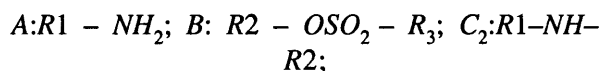
operating policy for the next experiment is determined based on the direction of improvement suggested by the chosen model. The results obtained are incorporated into the data set and another modelling cycle is done until model(s)/process mismatches are reduced enough and no further exploration is allowed. To exemplify the issues discussed, the following example is introduced.

3.4 Example 1

The following example has been adapted from Fotopoulos¹⁶. Consider the production of a drug intermediate C_2 , which is formed through the following mechanism:



C_1 , D_1 and E_1 are intermediates that cannot be measured; C_2 is the desired product, while D_2 and E_2 are undesired impurities. The symbols correspond to the following compounds:



Radical groups $R1$, $R2$ and $R3$ have had to be kept in secret.

The initial data set is made up of eight experiments as shown in Table I. The final concentrations of the measured species include a "biased noise" representing imperfect temperature control. The batch time for each run was 400 minutes.

The performance of the process is assessed using an economic-based index that attempts to maximize the amount of the desired product C_2 , penalize the use of raw materials and the formation

Table I(a)
Initial Data Set for Example 1. Operating Conditions

Run #	Batch time (min)	React. Temp. °K	Init. Amt. A(ml)	Init. Amt. B(ml)
1	400	385	360	200
2	400	385	360	300
3	400	385	600	200
4	400	385	600	300
5	400	410	360	200
6	400	410	360	300
7	400	410	600	200
8	400	410	600	300

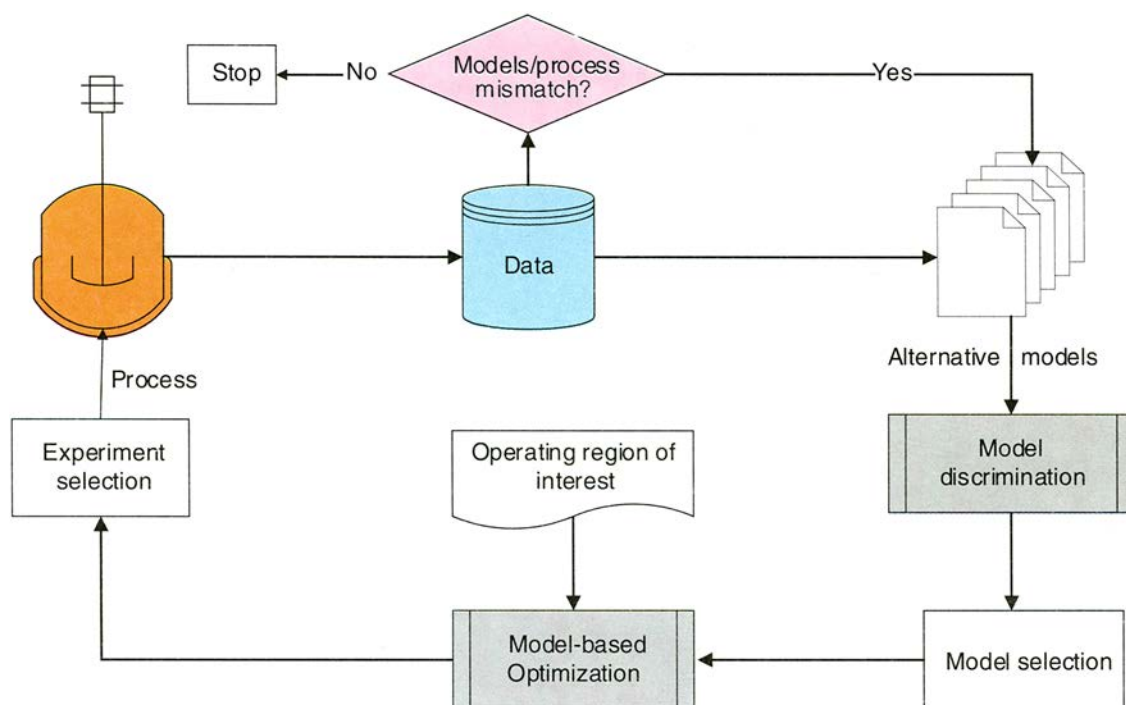


Fig. 4 Modelling for optimization cycle

Table I (b)
Initial Data Set for Example 1. Final Concentrations and Values of the Performance Index

Run #	Conc. Final A [mol/lit]	Conc. Final B [mol/lit]	Conc. Final C_2^f [mol/lit]	Conc. Final D_2 [mol/lit]	Conc. Final E_2 [mol/lit]	$J, \$/\text{day/liter}$
1	0.9939	0.0000	0.7587	0.6756	0.5586	+3.46
2	0.1627	0.1715	1.0623	0.9553	0.7690	+12.30
3	3.0122	0.0031	0.7594	0.6833	0.5588	-4.6
4	1.9941	0.0000	1.1394	1.0320	0.8400	+6.53
5	1.0122	0.0031	0.3865	0.4424	1.1789	-4.23
6	0.0657	0.0744	0.5615	0.6436	1.7000	+1.92
7	2.9956	0.0055	0.3948	0.4346	1.1922	-12.70
8	1.9938	0.0094	0.5641	0.6537	1.7588	-5.70

of unwanted by-products and minimize the batch time. The performance index is as:

$$J = \alpha_c \frac{(aC_2^f - bA^0 - cB^0) - \omega(A^0 + B^0 + D_2^f + E_2^f)}{t_b + t_s}$$

where:

A^0, B^0 , are the initial concentrations of reactants.,

$C_2^f, D_2^f, E_2^f, \dots$ stand for species concentrations at the end of the batch,

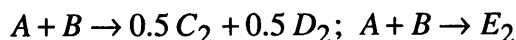
t_b is the batch cycle time in minutes and t_s is the reactor set-up time

α_c is a conversion factor to express J in $\$/\text{day/litre}$

a, b, c, d , are unit costs or values in $\$$

ω , separation costs.

As a result, the following approximate stoichiometry for the reaction system has been identified using singular value decomposition and Target Factor Analysis:



Using Power-law kinetics and assuming the rate orders were equal to the corresponding stoichiometric coefficients, three different models were obtained using different data sets. The estimated parameters were the pre-exponential factors and activation energies for the two reactions which are shown in Table II. The model **M0** is obtained using process data from the first eight runs, whereas **M1** is the model obtained using the runs # 1 through 9. The model **S1** in turn is obtained using process data for operating conditions corresponding to runs

Table II
Parameterizations of Alternative Process Models

Parameter	Model <i>M0</i>	Model <i>M1</i>	Model <i>S1</i>
A_1	1.43	40.91	2097.52
E_1	3798.0	6393.8	9900.3
A_2	1.105×10^9	5.462×10^9	3.530×10^{12}
E_2	20124.3	21317.3	26921.4

1 through 8 plus run #10. The operating conditions for run # 9 is the optimal solution predicted by *M0* whereas run # 10 is a "suboptimal" operating condition, yet in the gradient direction predicted by *M0*. Runs # 11 and #12 are the optimal operating conditions as predicted by models *M1* and *S1*, respectively. Row #13 in Table I is the actual optimum of the process.

To illustrate the impact of modelling errors in the modelling for optimization cycle, the possibility of following alternative directions of improvement provided by competing models are briefly discussed. Model *M0* predicts that the process performance will be significantly improved if the next experiment is run at its optimal operating policy. A value of 68.14 \$/day/litre is predicted for the performance index. When the experiment is actually done, a disappointing value of $J = 6.34 \text{ \$}/\text{day}/\text{litre}$ is obtained. The additional data (run # 9) did not invalidate model *M0* stoichiometry but parameters are changed which produces model *M1*. Assume that it is reasonable to believe that *M0* still provides a valid direction for improvement but the model/process mismatch may be due to extrapolation to too far different operating conditions. To test this hypothesis, a sub-optimal policy (run #10), yet in the gradient direction of *M0* is carried out. The resulting performance is $J = 23.17 \text{ \$}/\text{day}/\text{litre}$ which is only a bit lower than the predicted performance index of $J = 31.60 \text{ \$}/\text{day}/\text{litre}$. When the outcome of run # 10 is used to update model *M0*, model *S1* is obtained. There are now two possibilities for running the next experiment.

Table III summarizes the predicted and resulting outcomes for the four optimization runs. It seems that there exists a difference between the directions for improvement provided by alternative models. However, in the face of noise how can such differences be assessed?

4 Model Discrimination

The use of imperfect models as a guideline for evolutionary optimization of reaction systems can be understood as a search in the hypothesis space defined by model-based predictions of the process optimum. Each alternative process model provides a different hypothesis for an improvement direction and optimum location. Since hypotheses cannot be tested all at once, a key question to be answered is how the information from competing models is best used for deciding which operating conditions should be used for the next experiment? Also, how can the alternative model hypotheses be ranked or even compared? Answering these questions asks for developing model discrimination/selection criteria from the specific point of view of modelling-for-optimization. Also, a key issue to be addressed is that process data is noisy and biased which prevents the use of optimization techniques that are designed to deal with numbers rather than process outputs.

4.1 Correlation and Optimality

The outcome of each experiment provides sampled data for a performance index y which is a noisy function of a vector x of real-valued inputs defined over a feasible domain Ω that characterizes the controllable operating conditions for each experiment:

$$y = g(x) + \text{noise}, \quad x \in \Omega \quad \dots(1)$$

Given the constrained space of feasible operating conditions, the optimization task is to find an input vector that is close enough to the global optimum x^* of $g(x)$, using a rather short sequence of costly experiments. The main difficulty of dealing with noisy process outputs as defined in eq. (1) is making

Table III
Modelling Errors in the Modelling-for-Optimization Cycle

Run #	Concept	Batch time (min)	Temp. ^o K	Init. Amt. A (ml)	Init. Amt. B(ml)	J , obtained	J , predicted
9	<i>M0</i> optimum	81	368	500	411	+6.34	+68.14
10	Sub-optimal	240	383	490	331	+23.17	+31.60
11	<i>M1</i> optimum	92	370	510	406	+16.15	+56.00
12	<i>S1</i> optimum	132	380	506	420	+31.40	40.29
***	Actual optimum	172	377.1	491.2	437	+38.10	***

inferences about the optimum location for performance index $g(\mathbf{x})$ from a limited number of observations where the presence of outliers cannot be ruled out. Model-based optimization algorithms are often designed upon local gradients, which render them extremely vulnerable to noise and liable to “get lost” since they are not able to find a way to steadily progress towards an optimum.

Once a model has been selected to plan an optimizing experiment, the next crucial question is how to assess the actual capability of an imperfect model to locate the optimum operating condition in the face of noise and outliers. As a robust guideline for answering these questions on statistical grounds the following *monotonicity* assumption is proposed here: “the value of the chosen objective function $y(\mathbf{x})$ should exhibit evidence, in a statistical sense, of a definitive improvement as the operating conditions $\mathbf{x}$ move closer to the predicted optimum $\mathbf{x}^*$.” In mathematical terms this assumption states the existence of a monotonic relationship φ defined as follows³¹:

$$y(\mathbf{x}) = \varphi(\|\mathbf{x} - \mathbf{x}^*\|), \quad \varphi \text{ is monotonic} \quad \dots(2)$$

The observed data will fit this guideline of a local optimum to varying degrees, so one would like to quantify the degree of fit. A typical nonparametric measure used for this purpose is *Kendall's correlation coefficient* τ , otherwise known as Kendall's *tau*³². Typically, we wish to calculate τ for a candidate process model with $\mathbf{x}^* = \mathbf{x}_1^*$. To do this, we acquire the paired data $\|\mathbf{x}_i - \mathbf{x}^*\|$ and $y(\mathbf{x}_i)$, and denote their ranks by r_i and s_i . Next, arrange the r_i in ascending order. If $\mathbf{x}^*$ is a “promising” candidate for a minimum, then there should exist a strong positive monotonic association between these ranks (a negative correlation applies for a maximum). To measure this, we scored each paired difference $s_j - s_i, i = 1, 2, \dots, n-1$ and $j > i$ as +1 if the difference is positive and -1 if the difference is negative. Denoting the sums of positive and negative differences by n_c and n_d respectively, the equation for the Kendall's tau is:

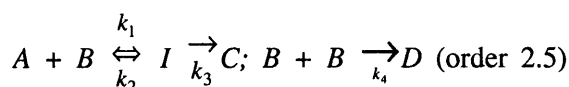
$$\tau_{\mathbf{x}^*} = (n_d - n_c) / \frac{1}{2} n(n-1) \quad \dots(3)$$

If the monotonic association between $\|\mathbf{x}_i - \mathbf{x}^*\|$ and $y(\mathbf{x}_i)$ is strongly positive, $\tau_{\mathbf{x}^*} \rightarrow 1$, there exists

conclusive evidence that the candidate model is defining an improvement direction pointing to a local minimum $\mathbf{x}^*$, whereas as $\tau_{\mathbf{x}^*}$ is decreasing, the hypothetical minimum $\mathbf{x}^*$ has less statistical support for its claim. For a local maximum, there will exist a negative correlation between $\|\mathbf{x}_i - \mathbf{x}^*\|$ and $y(\mathbf{x}_i)$, i.e. $\tau_{\mathbf{x}^*} \rightarrow -1$. Unless otherwise stated, let's assume we are seeking to maximize a process performance index. If the rankings of $\|\mathbf{x}_i - \mathbf{x}^*\|$ and $y(\mathbf{x}_i)$ are almost independent we can expect $\tau_{\mathbf{x}^*} \approx 0$, which highlights the poor performance of the process model as a guideline for improvement. The following simple example illustrates these ideas quantitatively.

4.2 Example 2

Consider a chemical reaction system conducted in an isothermal semi-batch reactor, which behaves according to the ‘unknown’ mechanism:



At a given operating temperature, the ‘true’ kinetic parameters (which are also assumed unknown) have the following *mean* values: $k_1 = 1.1355, k_2 = 5.0, k_3 = 4.5239$ and $k_4 = 3.5880$, which correspond to ‘perfect’ temperature control. In industrial-size batch reactors, temperature control is often far from perfect because of severe nonlinear behaviour and poor modelling. As a result, the final state of each run will show unsystematic variations (see Table IV). The reaction system is operated in a semi-batch mode where a stream of pure B ($[B]_{\text{feed}} = 0.2$ moles per litre) is added to a 1000

Table IV
Sampled Data for Example 1

Run #	x_1	x_2	t_{final}	J
1	45.0	7.25	102	0.5200
2	50.0	10.00	152	0.4732
3	48.0	8.00	118	0.5099
4	55.0	9.00	156	0.4711
5	38.0	12.00	125	0.5480
6	40.0	9.5	109	0.5221
7	44.0	11.00	139	0.4816
8	54.0	9.20	156	0.4660
9	58.0	8.60	162	0.4435
10	50.0	7.00	112	0.5299
11	52.0	9.60	155	0.4795
12	55.9	9.00	160	0.4400

litres vessel which initially contains 0.2 moles per litre of A and no B , and is filled to 50%. The objective is to maximize a productivity index (see Table I) for the process defined as follows:

$$J(x_f) = (1000 \times ([C]_f \times V_f)) / t_f \quad \dots(4)$$

For *safety* reasons in downstream processing, the final concentration of unreacted B cannot be greater than 0.0032 moles per liter. Hence, if the remaining B has a concentration greater than this maximum the batch must be allowed to continue further until a final time t_f where this threshold is achieved; as a result, the reactor productivity is decreased. Due to a heat extraction constraint the feed addition rate is limited to a maximum value of 12.0 litres min^{-1} , whereas because of production scheduling it is expected that the overall reactor cycle does not exceed 180 min.

Suppose data in Table IV has been obtained experimentally and the optimal solution x^* is hypothesized to be one with a semi-batch period (x_1) of 40 min. and a feed rate (x_2) of 9.5 litre/min. In order to assess for this optimum's degree of fit, let's test for a positive monotonic association

between $y(x)$ and $\|x - x^*\|$. This relationship is illustrated in Fig. 5. To test for the hypothesis that the optimum is $x = (40, 9.5)^T$, the value of t_K of and its significance needs to be calculated. Table V provides ranked data from Fig. 5; the paired ranks are used to calculate the statistic for the hypothetical optimum yielding $t_K = -0.4545$ with $n_c = 17$ concordances and $n_d = 49$ discordances.

Testing for the significance of the hypothesis $H_0: \tau = 0$ against the alternative $H_1: \tau < 0$ requires a precalculated value of the Kendall's τ critical values for $n=12$ and a chosen significance level $(1-\alpha)\%$ ³³. Statistical tables give nominal 5 and 1 per cent critical values for significance when $n=12$ in one-tail test as -0.3929 and -0.5455. Corresponding values for $n_d - n_c$ are 26 and 36, whereas an approximated Monte Carlo estimate of the exact one-tail probability P of rejecting H_0 when in fact it is true is 0.0396. Since, we reject H_0 and accept the postulated optimum with a confidence of

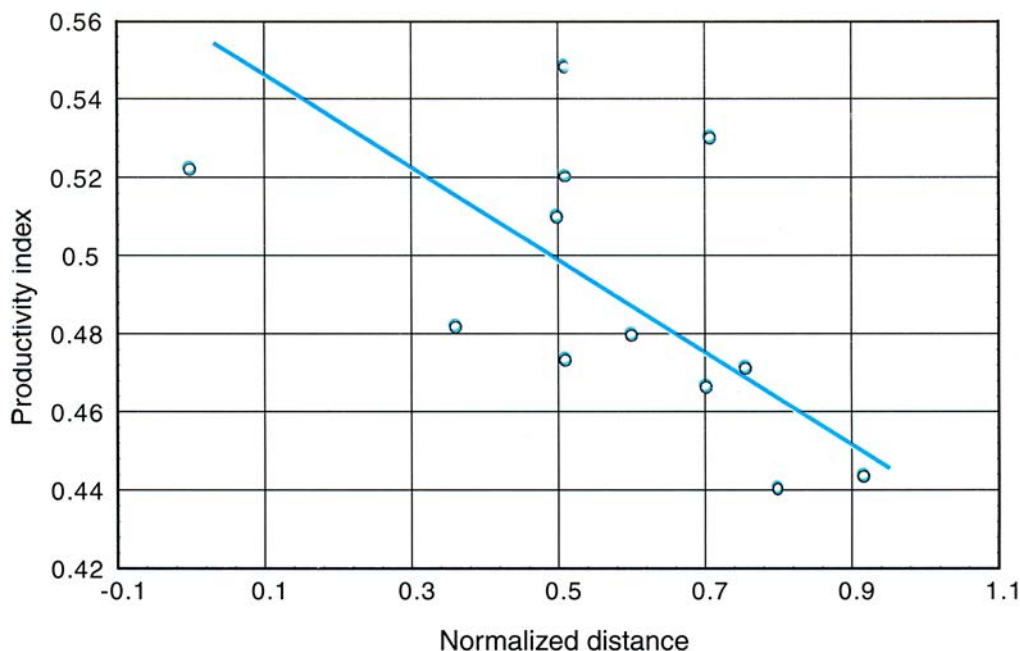


Fig. 5 Correlation and local optimality for the data set in Example 2

Table V
Ranked Data for Example 1

x -rank	1	2	3	4	5	6	7	8	9	10	11	12
y -rank	10	7	8	5	12	9	6	3	11	4	1	2

95%, on the other hand the available data does not provide strong statistical support to accept the postulated optimum at $x = (40, 9.5)^T$ with a confidence level of 99%.

4.3 Concordance and Rivalry

At first glance it may seem reasonable to rank alternative models using the Kendall's tau. Table VI compares alternative models in Example 1. Since the monotonic relationship (2) is based on local considerations, only run # 9 through #12 are used for Kendall's tau calculation. As can be seen only *S1* does provide enough statistical evidence to claim that its predicted maximum is credible enough. Biased and noisy data ask for a more elaborate measure of concordance among alternative models. For this, an extension of the foregoing analysis using a pair-wise bisection approach is proposed here. The core idea is as follows: if a perpendicular bisector is created between a pair of experimental points x_1 and x_2 in the dataset, two process models are *concordant* if their predicted

Table VI

Kendall's Correlation Coefficient for Alternative Optima

Postulated optimum, x^*	Kendall's tau
M0 (run # 9)	+ 0.61
M1 (run # 11)	+ 0.40
S1 (run # 12)	- 0.41
Actual optimum	- 0.60

maxima are both on the same side of the bisector as the point providing the greatest $y(x_i)$. It is easy to imagine creating bisectors for all possible $(n-1)$ pairs in the data set and counting the number of bisectors that satisfy the model concordance

criterion. The number of satisfied bisectors minus the number of unsatisfied bisectors, divided by the total number of bisectors, will have the same distribution as the Kendall's tau.

Fig. 6 contains an example of the proposed pair-wise bisector approach for Kendall's tau and its calculation with three data points and two tendency models and two input dimensions. Each circle represents an actual experiment with the objective value inside the circle whereas the coordinates are two independent inputs describing the operating condition. Solid lines are the perpendicular bisectors of each pair of points. For each pair's bisector, a solid triangle indicates the direction of the greater of two points. A +1 indicates whether both predicted optima x_1^* , x_2^* are on the same side of the bisector as the greater of the two points (i.e., a concordance). A -1 indicates the opposite, i.e. at least one of the predicted optima is on the other side of the bisection from that of the larger of the two points.

When the notion of pair-wise bisection is integrated with the monotonicity assumption stated above (see eq. (2)), a straightforward procedure to differentiate between concordant and rival models is: **Step 1.** For model "1" use its predicted optimum x_1^* along with the data set to rank the values of the objective function $y(x_i)$ in ascending order of $\|x_i - x_1^*\|$. Let's denote θ_i^1 the set of ranks for the objective function thus obtained.

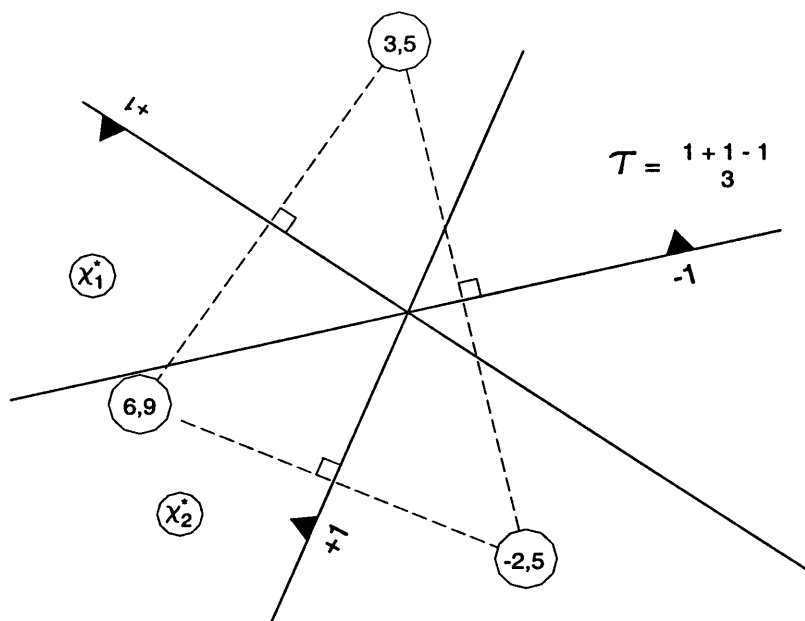


Fig. 6 The bisection approach to model concordance

Step 2. Repeat Step 1 for model “2” using its predicted optimum x_2^* ; denote θ_i^2 the resulting set of ranks for the objective function.

Step 3. Apply the Kendall’s tau procedure to test for independence between θ_i^1 – ranks and θ_i^2 – ranks, against the alternative hypothesis for positive or negative correlation between ranks for the chosen degree of confidence α .

If the hypothesis for positive correlation between θ_i^1 – ranks and θ_i^2 – ranks can be accepted, these models can be considered as concordant to the chosen level of significance. Conversely, if the hypothesis for negative correlation between ranks is the one to be accepted with a level of confidence, tendency models are considered as rivals to this extent. Finally, if the hypothesis for no correlation is accepted, the available evidence is inconclusive about concordance or rivalry between models.

4.4 Example 1 (cont’d)

In Table VIIa the ranks obtained for models *MI* and *SI* after applying the three steps above are given, whereas Table VIIb depicts the ranks corresponding to *MI* and the actual process. Using the ranks of Table VIa, the Kendall’s tau is $\tau = -1$ which indicates that, according to the information available, models *MI* and *SI* are not concordant but rivals. Conversely, when *SI* is compared to a “perfect” model (i.e., the one that predicts the optimal solution) the resulting Kendall’s tau is $\tau = +0.8$, a strong indication of concordance. It can be concluded from the available data set that *SI* is the model with better chances of providing a reliable estimation of the process optimum.

5 Model-Based Evolutionary Optimization

Following model discrimination, the evolutionary search towards the process optimum poses some important questions. The first one is: how to value

the capability of a model to predict the process optimum? In modelling for optimization the role of modelling is reducing the uncertainty about the optimum location. Accordingly, the values of alternative models are somewhat related to their capability to provide useful guidelines for improvement. Assuming each model’s estimated value is known, how can this knowledge be systematically used for model selection while balancing exploitation with exploration? Given that models are inherently imperfect and run outcomes are scarce and noisy, the design of a model selection procedure must resort to active learning from process data. Finally, once a model has been selected, how is the operating policy for the next experiment decided?

5.1 Model Valuation

The amount of improvement that follows each model selection gives some information about how good the decision was, but it says nothing at all about whether the model hypotheses were correct or incorrect, that is, about whether or not choosing a given model to define an improvement direction is the best decision. Correctness is a relative property of model predictions that can be determined only by trying them all and comparing the actual improvement that has been obtained. Hence, the problem of model selection is inherently one requiring explicit search among the alternative model hypotheses. Model selection requires some form of generate-and-test method whereby model predictions are valued, operating policies are tried, the outcomes are observed, and selectively the most effective model hypotheses are retained. This is learning by evaluative feedback, which is the basis for methods of evolutionary optimization.

A key issue for model selection is assigning *a priori* a value to alternative model predictions based on the cumulated experience. The estimated value of a model is the expected reward given that model is selected for defining the operating conditions in the next run. In the face of modelling errors the value of a model is mainly a trade-off between improvement prediction and the degree of uncertainty about that prediction. The less uncertain is the predicted gain in performance, the more valuable is the model prediction. For a given model structure, model valuation requires to assess the impact of parametric uncertainty on the uncertainty of the process optimum location. The degree of

Table VII(a)

MI-ranks Against SI-ranks

<i>MI-ranks</i>	2	1	4	5	3
<i>SI-ranks</i>	4	5	2	1	3

Table VII(b)

SI-ranks Against Optimal Ranks

<i>SI-ranks</i>	4	5	2	1	3
Optimal ranks	5	4	2	1	3

confidence in a predicted optimum can be measured using a “covariance matrix” for the optimal policy $u^*(t)$, C_u^* , defined as follows²⁰:

$$C_u^* = \partial u^* (\partial u^*)^T = \left(\frac{\partial u^*}{\partial p} \cdot \partial p \right) \left(\frac{\partial u^*}{\partial p} \cdot \partial p \right)^T \quad \dots(5)$$

$$= \left(\frac{\partial u^*}{\partial p} \right) \partial p \partial p^T \left(\frac{\partial u^*}{\partial p} \right)^T = \Xi C_p \Xi^T$$

where p is the vector of model parameters, $C_p = \partial p \partial p^T$ is the parameter covariance matrix and $\Xi = \left(\frac{\partial u^*}{\partial p} \right)$ is the sensitivity matrix. The matrix C_u^* can be used to narrow down the uncertainty about the location of the process optimum through confidence limits for each j th policy parameter defined as:

$$\Delta u_j = \pm (C_u^*)_{jj}^{0.5} t_{1-\alpha/2, v} \quad \dots(6)$$

where $t_{1-\alpha/2, v}$ is the value of the two-sided t -distribution for $v = n_d - n_u$ degrees of freedom (n_d is the number of experimental points and n_u is the number of policy parameters) and $100(1-\alpha)\%$ is the confidence of locating the optimum policy.

The approximate $100(1-\alpha)\%$ confidence region around the optimal operating condition u^* is the hyper-ellipsoid defined by

$$(u - u^*)^T C_u^{-1} (u - u^*) \leq \chi_{dg}^2(\alpha) \quad \dots(7)$$

where χ^2 is the chi-squared distribution³⁴ and dg is the number of degrees of freedom (number of experimental points minus the number of model parameters). The eigenvectors of C_u give the direction and the eigenvalues give the length of the axis of the hyperellipsoid defined in eq. (7). It is worth noting that the volume of this hyper-ellipsoid is inversely proportional to the determinant of C_u^{-1} , $Det(C_u^{-1})$.

On the basis of the above analysis, the value $Q(u^*)$ of a model in determining an optimal operating condition over the region of interest is defined here as the quotient between the predicted

improvement $\Delta \hat{J}$, , and a measure of the uncertainty in model predictions, e.g. the volume of the hyper-ellipsoid in eq. (7). In mathematical terms, the estimated value of a model in predicting the optimum location at u^* can be conveniently written as

$$Q(u^*) = \Delta \hat{J}(u^*) \cdot Det(C_u^{-1}) \quad \dots(8)$$

where the improvement $\Delta \hat{J} = \hat{J}(u^*) - J(u_{\max})$ is calculated with regard to the operating policy $u_{\max}$ that has the greatest value of J in the current data set.

This approach to model valuation is somewhat related to the notion of *D-optimality* in optimal experimental design³⁵ which carefully selects operating conditions in the sequence of runs so as to achieve a reduction in the uncertainty of model estimated parameters (see later for details). However, in modelling for optimization the operating conditions must be chosen to reduce the uncertainty in the optimum location rather than model improvement. Hopefully, as the experimental design strategy is increasingly biased in the gradient direction given by the highest-valued model, the uncertainty about the process optimum is decreased.

5.2 Model Selection and Active Learning

The information content of the process data generated in the modelling for optimization cycle has a definitive impact on the rate of learning and the degree of improvement. The appropriateness of data depends on the performance function that is being learned and intrinsic process variability. A sequential experimental design strategy for determining the most profitable operating conditions can be designed to follow either a *greedy* or an *active learning* approach^{7, 36, 37}. In the first case, the sequence of experimental trials is simply the result of attempting to maximize the *exploitation* of current knowledge derived from alternative model hypotheses (structure + parameters), including estimated optimum locations, estimated confidence regions and improvement directions. No attempt is thus made to influence the generation of data in order to circumvent modelling errors and poor extrapolation by carrying out experiments that may help reducing the uncertainty regarding the location of the optimum. Active learning is a form of learning that assumes that the learner has control over what part of the input domain it receives

information about. Thus, experiments are pursued to systematically reduce the uncertainty regarding optimum location, i.e. maximize the value of a model.

If the actual value of each model is known, then the model selection problem would be trivial to solve: the model with highest value must be always selected. We call this a *greedy* model selection which seeks to exploit current knowledge about the values of alternative models. Exploitation is the right thing to do as long as the uncertainty about the optimum operating policy predicted by each model has been reduced enough. However, if the uncertainty is such that the optimum locations predicted by competing process models might actually be better than active learning is called for. Active learning is understood here as a form of learning where modelling for optimization explicitly allows some room to follow models whose predicted optima exhibit higher levels of uncertainty. The experimental design strategy is thus deliberately designed to allow some degree of exploration, which is essential to deal with modelling errors and scarce data. The problem with exploration is that it is costly and, sometimes, can also be prohibitively sub-optimal or even produce wasted batches.

The balance between exploration and exploitation in model selection can be achieved in different ways. One alternative is the ℓ -greedy approach, which consists of selecting the “best” model most of the time, but every once in a while, say with small probability ℓ , explore apparently less promising models. This means choosing operating conditions for which there currently exist weaker or discouraging statistical evidence that an optimum can be found nearby. The main drawback of this method is that when exploring it chooses equally among all the candidate models without any account for their values. A better method is to vary the model selection probabilities as a graded function of the size of the confidence region around the predicted optimum. The predicted optimum exhibiting the highest degree of confidence according to the available data (the greedy model) is still given the maximum selection probability, but the set of all hypotheses (models) is valued according to their parametric uncertainty and concordance or rivalry with the greedy model. This is called the *softmax* criterion where the probability of selecting the j th model for defining the operating condition for the next run is defined as:

$$e^{Q_j/T} / \sum_{i=0}^r e^{Q_i/T} \quad \dots(9)$$

where r is the number of candidate models to choose from (including the greedy one), $Q_i = 1, \dots, r$ are their corresponding *values* (calculated as explained above) and T is a positive parameter called the “temperature”. High temperatures cause all the r candidates to be nearly equiprobable. As the temperature is lowered, policies with higher values have greater chances of being selected. In the limit as $T \rightarrow 0$, the bias towards pure exploitation is total, i.e. the probability of selecting the greedy model tends to 1, and no room for exploration is left.

The knowledge about rivalry or concordance among alternative models with regards to the model with highest value can be used to advantage. One algorithm that accounts for this is called ε -*Softmax*. In this case, with probability $(1 - \varepsilon)$ the model with highest estimated value or one that is concordant with it is chosen; with small probability ε one of the rival models is chosen. The ε -*Softmax* algorithm for model selection is summarized in Fig. 7. It is worth noting that after having decided which group of models to use (concordant or rival) model

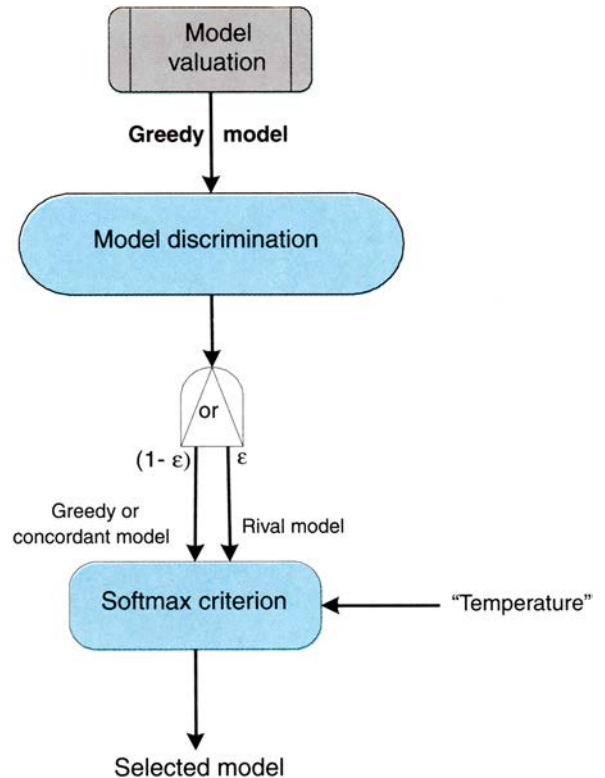


Fig. 7 The ε -Softmax procedure to model selection

selection is still based on eq. (9), i.e. the softmax criterion.

An important consideration in the application of the *Softmax* and ε -*Softmax* algorithms is the cooling procedure chosen to gradually diminish the amount of exploration. One simple alternative is to practice a linear “cooling” procedure as follows. Each time a new operating condition is actually tried, and the results are added to the experimental data set, the temperature parameter T^{iter} is decreased using:

$$T^{iter} = \xi T^{iter-1}, \quad 0 < \xi < 1 \quad \dots(10)$$

where ξ is the rate of cooling. More elaborated cooling strategies, like exponential cooling, can also be used to advantage as it is shown later.

5.3 Experimental Planning

Following model selection, the operating policy for the next experimental run must be decided. Because of the uncertainty in model predictions, it may look a bit risky to run the next experiment at the predicted optimal operating policy, mainly if this optimization run leads the process far from operating conditions in the current data set. Instead a more cautious exploration along the gradient direction would be preferable. The chosen model provides a direction for improvement which can be defined as:

$$u_{n+1} = u_{\max} + \lambda (\hat{u}_{n+1} - u_{\max}), \quad 0 \leq \lambda \leq 1 \quad \dots(11)$$

where $u_{\max}$ is the process optimum in the current data set, i.e. the best operating policy found so far; $\hat{u}_{\max}$ is the optimum location as predicted by the selected model. The uncertainty for a sub-optimal operating policy SO defined along the improvement direction in eq. (11) can be calculated from:

$$\begin{aligned} \Xi_{SO} &= \frac{\partial}{\partial p} [u_{\max} + \lambda (\hat{u}_{n+1} - u_{\max})] \\ &= \lambda \cdot \frac{\partial \hat{u}_{n+1}}{\partial p} = \lambda \Xi, \quad 0 < \lambda < 1 \quad \dots(12) \end{aligned}$$

This new value of Ξ_{SO} can be used as before (see eq. (5)) to calculate a covariance matrix for a sub-optimal policy SO :

$$C_u^{SO} = \lambda^2 \Xi C_p \Xi \quad \dots(13)$$

Thus, the degree of confidence in the suboptimal policy prediction can conveniently be increased as the parameter λ is lowered. The amount of improvement that can be gained will also be

reduced, however. There can be many alternative ways to trade-off expected improvement with uncertainty. The best approach is seeking for the maximum value along the line in eq. (11). This is explained graphically in Fig. 8. An easier alternative is allowing a level of uncertainty for the next run that is not greater than the uncertainty associated with the model with highest value. Alternatively, the amount of extrapolation λ can be profiled down along with the “temperature” parameter T used above to limit the room for exploration.

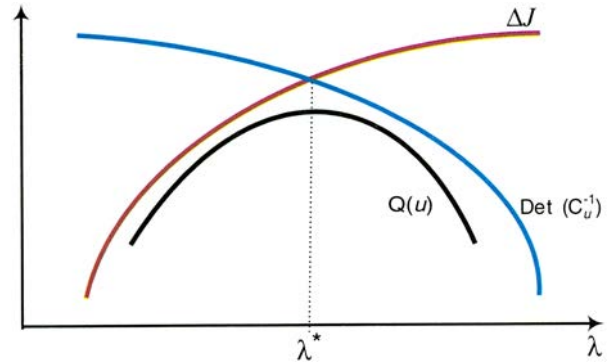


Fig. 8 The trade-off between the predicted improvement and uncertainty

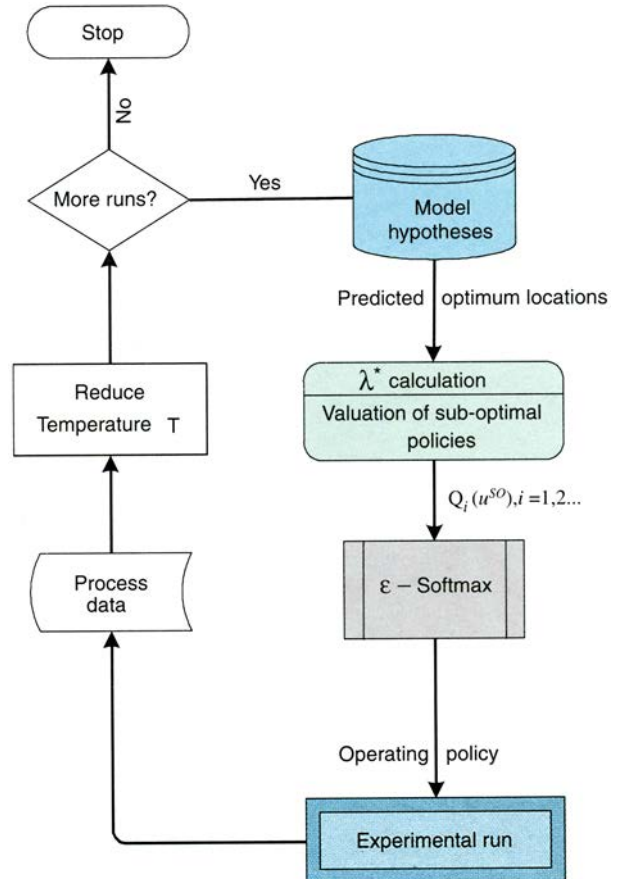


Fig. 9 Model-based evolutionary optimization algorithm

The algorithm for model-based evolutionary optimization is summarized in Fig. 9. In a given iteration, the operating policy corresponding to λ^* is first determined for each alternative model. Each model is assigned the value corresponding to the sub-optimal policy defined by its λ^* . Alternative models are then tested for rivalry or concordance regarding the model with highest value. Later on, the *Softmax* criterion is applied to decide which operating conditions to use in the next run. The experiment is done. Once the experiment is done the results are incorporated into the data set and the “temperature” is lowered. The hypotheses about model structures and parameters are updated and a new iteration is completed. The algorithm stops when there is no more room (time and/or money) for modelling runs. The operating policy with highest value is chosen as the process optimum.

5.4 Example 1 (Continued)

The model-based evolutionary optimization algorithm has been applied to the Example 1 discussed earlier. A linear cooling profile is used to diminish exploration as modelling runs are added to the initial data set (Table I). The rate of cooling is chosen such that after five new modelling runs exploration is over. The results obtained are summarized in Fig. 10. The remaining improvement gap from run # 13 on is mainly due to modelling errors and noisy outputs which cannot be reduced further by simply doing more experiments.

6 Batch End-Point Constraints

Imperfect and incomplete understanding of process behaviour compounded with intra- and inter-run variations not only prevents achieving an optimal operation of batch processes, but also gives rise to potential risks of violating product end-use properties, ecological or safety constraints. This section discusses a sequential experiment design strategy based on active learning with an imperfect model to accomplish the specific goal of *modelling for optimization* under batch end-point constraints.

6.1 Problem Statement

Optimal operation of batch processes is typically based on a performance function $J(x_f)$ that involves the final state x_f of the batch run, which in turn also needs to satisfy a set of endpoint constraints $g(x_f) \geq 0$. For a given operating policy u , unknown initial conditions and uncontrollable intra-run variations make batch run outcome x_f uncertain and very difficult to model accurately. Even under laboratory conditions, end-point reproducibility is as low as 5%⁴. Significant variability of each run outcome does make it difficult to measure the actual quality of using a given policy u in terms of the chosen objective function. Also, outcome variance gives rise the problem of assessing the probability of satisfying endpoint constraints for each operating policy defined over the feasible region of policy parameters. Bearing in mind these considerations,

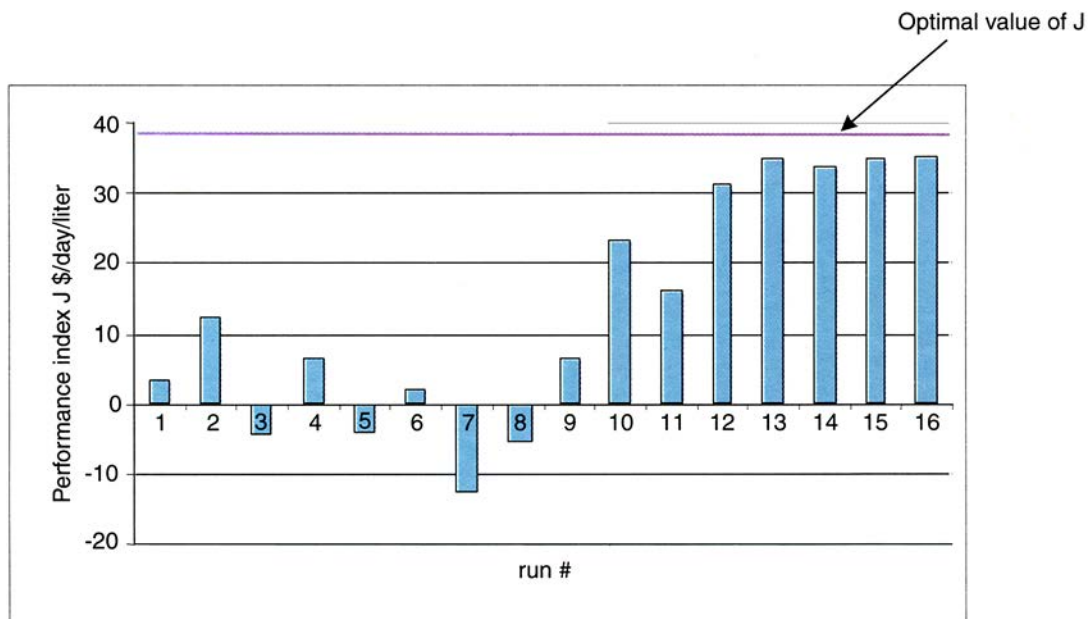


Fig. 10 The learning curve for Example 1

the problem of *learning* improved operation using outcomes of successive runs consists of finding an operating policy that incrementally approximates the greedy policy u^* , which solves:

$$\underset{u}{\text{Optimize}} \ E[J(x_f)] \quad \dots(14)$$

where: $x_f = \chi(x_0, u) + \text{noise}$ (known through experiments), and

Subject to:

$$g(x_f) \geq 0, \quad \text{Endpoint constraints,}$$

$$u_i^L \leq u_i \leq u_i^U, u_i \in u \text{ Policy parameters bounds}$$

where $E[\bullet]$ is the mean value operator, and J is a figure of merit for the final state that may involve not only the expected performance but also its variability. The mapping χ represents an *a priori* unknown relationship describing the averaged effect of each operating policy and batch initial condition on the final state. Thus, the actual influence of a given policy on x_f progressively unfolds as more replications for the operating policy are carried out. Changes in x_0 are mainly the result of inter-run variations such as catalyst activity, reactant/solvent recycling or raw material quality. The term “noise” denotes unknown variations or disturbances acting on the process whilst it is being processed, e.g. poor temperature control or inadequate mixing. This noisy component of x_f is typically due to intra-run variations resulting in different process behaviours affecting outcome reproducibility in a biased way. Thus, the stochastic nature of “noise” is often non-Gaussian and skewed, which makes very difficult to use standard statistical techniques for plant data analysis.

From the standpoint of modelling for optimization two issues, feasibility and performance, exist that need to be integrated together when assessing the true *value* or *utility* of the solution found to the problem in eq. (1). For feasibility, the value of each operating policy should measure the probability of satisfying endpoint constraints³⁸. The probability of successfully completing a batch run is estimated here analytically using an approximated model. Statistical analysis of the uncertainty associated with model parameters is used to assess the risk of endpoint constraint violation. Cubature integration techniques³⁹ or other quadrature formulas

aimed at efficient sampling are readily available for doing this calculation. The maximum achievable value $\Pr_u[g(x_f) \geq 0]$ will depend on the magnitude of uncontrollable disturbances affecting each batch outcome, i.e. x_f . For performance, the value of each policy should account for the expected value of $J(x_f)$. The quality of experimental data generated whilst modelling for optimization is then measured in terms of the increasing capability of the resulting model to help assess the *value* of policies.

The model-based evolutionary optimization strategy proposed here systematically shrinks the Region of Interest **ROI** for each policy parameter $u_i, i = 1, \dots, r$, over which process optimization is constrained to take place. Initially, **ROI**₀ corresponds to the overall feasible region for defining a policy as defined in eq. (1). By cutting away unpromising regions within **ROI**₀, the algorithm is aimed at generating a sequence of **ROI**₁, **ROI**₂, ..., **ROI**_{iter}, **ROI**_{iter+1}, ..., such that

$$\mathbf{ROI}_{iter} \supseteq \mathbf{ROI}_{iter+1} \quad \dots(15)$$

Eventually, a final **ROI** is obtained for which all possible operating policies are “non-superior” to each other. For this subset of policies, different values for x_f can only be explained on the basis of the “noise” component. Thus, further improvements can result only from reducing the variance of “noise” and, measurements permitting, by operating policy improvement through on-line (intra-run) optimization that compensates for changes in the batch initial conditions x_0 .

6.2 Valuing Operating Policies

Modelling for optimization under uncertainty is understood here as *learning* from sampled reinforcements in a succession of experimental runs to selectively bias policy selection towards the most promising operating conditions. To assess retrospectively the goodness of previous decisions in this regard, a way is required to reward alternative operating policies in the face of uncertainty. Let's consider that each policy u has an expected endpoint *value* given that this policy has been tried infinitely. In a practical sense, this true value can be approximated by averaging the $J(x_f)$ actually obtained when a policy u is chosen, weighted by a model-based estimation of $\Pr_u[g(x_f) \geq 0]$.

Sampled estimations of values aim to approximate the “true” endpoint value $Q(u)$ for a given policy u defined as follows

$$Q(u) = \left(\frac{1}{r} \sum_{e=1}^{r \rightarrow \infty} J^e(x_f) \right) \times \Pr_{\mathbf{u}}[g(x_f) \geq \mathbf{0}] \dots (16)$$

The modelling for optimization cycle is aimed at incrementally improving the estimation of the endpoint value for each policy u previously tried so as to make a sound decision for choosing a handful of operating policies that are experimentally tried. Because of the uncertainty involved, value estimation creates the need for selective exploration during learning. In this section, exploration refers to either choosing again (replicating) one of the policies already tried or trying an entirely new one. In the latter case, deciding where over **ROI** exploration is deemed necessary demands the use of a function approximation technique (e.g., regression or neural network) to extrapolate/interpolate the values of control policies to unseen operating conditions. Hence, at any time there is at least one policy within a given **ROI** whose estimated value seems to be the greatest. We call this a *greedy* policy. If this policy is chosen for the next run, it can be said that we are “exploiting” our current knowledge about the values of operating policies.

Exploitation is the right thing to do assuming that policy value estimates are sufficiently good and process data are abundantly available all over **ROI**. However, if uncertainty is such that there would be other operating policies whose true values are in fact higher than the current greedy policy, then exploration is called for. Since it is not possible both to explore and to exploit with any single policy selection a conflict between these orthogonal objectives arises. To obtain more reward, we must prefer control policies that have been already tried and found to be effective in terms of the endpoint objective and constraints. But to discover such policies, we must try operating conditions that have not been selected before. Exploitation is needed to obtain more reward whilst exploration is needed to find better policies to exploit in the future. The dilemma is that neither exploitation nor exploration can be exclusively pursued whilst doing modelling for optimization without failing to satisfy the endpoint constraints. Having said this, the solution algorithm needs to be designed to selectively try

informative control policies, and progressively favours those that appear to be the best, based upon actual experience. Also, and because of uncertainty, each operating policy must be tried a minimum number of times (say 2 or 3) to gain a reliable estimate of its expected endpoint value.

The balance between exploration and exploitation can be achieved in different ways. One alternative is the ℓ – greedy approach, which consists of selecting the “best” policy over **ROI** most of the time, but every once in a while, say with small probability ℓ , select one of the previously tried policies. The main drawback of this method is that when it explores it chooses equally among all the policies already tried without any account for their endpoint values. A better method is to vary the policy selection probabilities as a graded function of estimated values. The policy having the highest value according to current knowledge is still given the maximum selection probability, but all previously tried policies are ranked according to their value estimates.

6.3 Sequential Experiment Design Strategy

Given as input a data set DS_0 from previous batch runs:

$$DS_0 = \{(u_1 \rightarrow Q(u_1), (u_2 \rightarrow Q(u_2) \dots (u_n \rightarrow Q(u_n)))\} \dots (17)$$

such that for any operating policy $u_j, j = 1, 2, \dots, n$, the

i -th parameter $u_i \in u$ satisfies the **ROI**₀ constraints.

The algorithm begins with the assignment $iter := 0$, and it then carries out iteratively the following steps (see Fig. 11).

Ranking

Each policy $u_j, j = 1, 2, \dots$, for data points included in DS_{iter} is given a figure of merit to identify where cutting away an unpromising, yet feasible sub-space of the policy parameters $u_i, i = 1, \dots, r$ is deemed as appropriate. At the onset of experimentation it is better to rank policies using an estimation of their probabilities of satisfying endpoint constraints. Doing this requires resorting to an approximated process model. As soon as modelling for optimization progresses such that policy value estimations are more reliable, ranking policies in DS_{iter} according

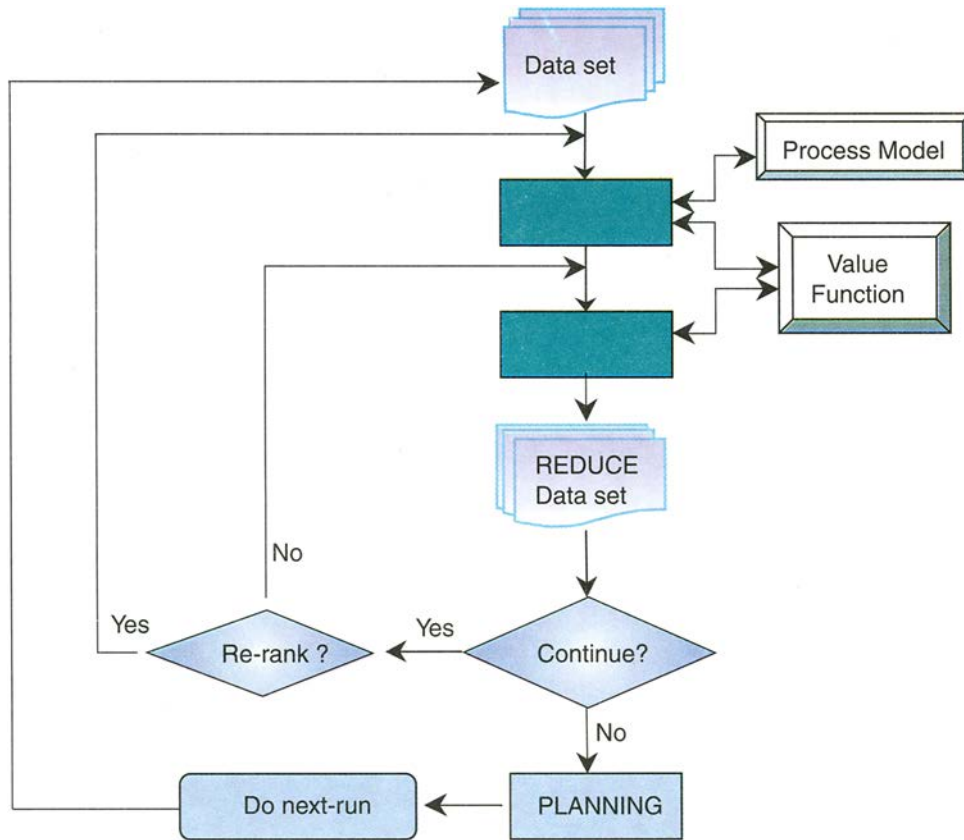


Fig. 11 Learning to satisfy batch end-point constraints

to their values is a much better approach. Once the ranking has been done, it is easy to identify the “less” promising policy $u_{\min}$ in the data set, and as such, it will be used to define a “cut” to reduce the size of ROI_{iter} . It is worth noting the decisive role that policy ranking will have in the sequence of decisions produced by the experimental design algorithm.

Cutting

Let $(u_{\min} \rightarrow Q(u_{\min}))$ be the data point that is predicted to be the “worst” within DS_{iter} . To define a cut of ROI_{iter} , all we need is an estimation of the direction of the steepest gradient $\nabla_u Q(u)$ at $u = u_{\min}$. This requires using some function approximation technique to extrapolate/interpolate the values of policies over ROI_{iter} . There might be several approaches to follow for this purpose, e.g. using a regression model, a neural network or a fuzzy system. Because of the scarcity of data, the use of a simple linear (in the parameters) regression model is considered enough here:

$$\hat{Q} = X^T \beta + \varepsilon \quad \dots(18)$$

where $X = X(u)$ is a $1 \times m$ matrix function of the control variables with rank m , β is an m -vector of unknown parameters and ε a normally distributed scalar variable with zero mean and variance σ^2 , is sufficient for our purpose. Moreover, since DS_{iter} often will consist of very few, weirdly distributed data points, simply fitting a quadratic is proposed here:

$$\hat{Q}_j(u_j) = c + b^T u_j + \frac{1}{2} u_j^T A u_j \quad \dots(19)$$

where c , b and A correspond to the set of fitting parameters to be obtained by the “least-squares” criterion. The local gradient at $u = u_{\min}$ is then easy to calculate as $\nabla Q = b + A u_{\min}$. A cut to DS_{iter} is made using the half-plane perpendicular to ∇Q so that:

$$ROI_{iter+1} = ROI_{iter} \cap \{u \mid (u - u_{\min}) \cdot \nabla Q \geq 0\} \quad \dots(20)$$

For the next cut, the data points to be considered are reduced to:

$$DS_{iter+1} = DS_{iter} - (u_{min} \rightarrow Q(u_{min})) \quad \dots(21)$$

and, optionally, a re-ranking of remaining data points in DS_{iter+1} is done; *Ranking/cutting* is continued this way until one of the two following conditions apply:

(i) There is not enough data present in DS_{iter+1} to determine all fitting parameters in c , b and A . The immediate result of insufficient data is matrix inversion problems;

(ii) Confidence interval around $Q(u^*)$, for the estimated optimum operating policy $u^* = -A^{-1}b$, exceeds $\|Q(u_{max}) - Q(u_{min})\|$ corresponding to best and worst points in DS_{iter+1} , respectively. Once no more cuts are allowed, proceed to the *planning* step below.

Planning

At this stage, the crucial decision of defining the policy u_{next} to be used in the next experiment needs to be made. As mentioned above, the key question is which criterion is better in order to gain the greatest amount of useful new information. The options are either to exploit or to explore further. From the point of view of exploitation, it seems attractive to consider the greedy decision of choosing the operating policy for the next run as $u_{next} = u^* = -A^{-1}b$. However, this excessive bias towards exploitation is unsound unless the information gathered in previous runs provides enough evidence to support that $\Pr_u[g(x_f) \geq 0]$ is high enough all over the reduced ROI and the confidence interval around u^* is sufficiently small. Otherwise, the next run operating policy should be decided on the grounds of seeking for a trade off between pure exploitation against further exploring the reduced ROI provided by the *cutting* stage above. From the point of view of exploration, we consider the question “where should one explore further over the reduced ROI in order to achieve the maximum reduction in the uncertainty about the location of the optimal policy?” To answer this, concepts of optimum experiment design (OED) theory are needed. According to the OED the next run policy should be chosen so as to minimize the variance of policy value estimations over ROI. To gain insight on how this can be attained, it is convenient to consider the variance-covariance

matrix of the least square estimate $\hat{\beta}$:

$$\text{Var}(\hat{\beta}) = (X^T X)^{-1} \sigma^2 \quad \dots(22)$$

For any policy $\hat{u}$ over the ROI, the estimated value is $\hat{Q} = X^T \hat{\beta}$ with a variance given as

$$\text{Var}(\hat{Q}) = X^T (X^T X)^{-1} X \sigma^2 \quad \dots(23)$$

It is worth noting from eq. (22) that the variability of the estimate $\hat{\beta}$ is influenced by the control policies that make up the matrix $(X^T X)^{-1}$. Furthermore, the quality of the regression model's prediction for the value of an operating policy $\hat{u}$ over ROI also depends on the matrix $(X^T X)^{-1}$. Accordingly, a good decision for the new operating policy to be explored could be any criterion that makes this matrix, in some sense, “small”. The most popular criterion is called *D-optimality*⁴⁰ and is based on choosing a new operating policy $u_{next} = u_D$ over the final ROI so as to minimize the determinant of $(X^T X)^{-1}$. *D-optimality* as a criterion for optimum experiment design has a number of desirable properties, e.g. minimum variance predictions and providing the smallest confidence ellipsoid for parameter estimates $\hat{\beta}$.

In addition to the control policies in DS_0 which can be tried again, there exist also two new optional policies to be used for the first time, namely u^* and u_D . The decision is taken here using the *softmax* criterion and value-dependent probabilities calculated as shown in eq. (9). The values of u^* and u_D are estimated using the simple model given in eq. (19). The overall modelling for optimization algorithm is summarized in Fig. 11. It is worth noting that determining the location of u_D is done here by means of a simple optimization problem as follow:

$$\max_{u_D} \text{Det} [(X^T X)^{-1}] \quad \dots(24)$$

Subject to: $u_D \in \text{ROI}_{final}$

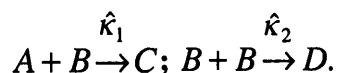
The decision variable u_D only arises in the $(m+1)$ -th dummy column added to X . Since the feasible

region for maximizing $\text{Det}(X^T X)^{-1}$ is constrained over the final **ROI** resulting from the cutting step the exploration is thus much less risky.

6.4 Example 2 (Continued)

The specifications on the final state of each batch are as follows. For *safety* reasons in downstream processing, the final concentration of unreacted *B* cannot be greater than 0.0035 moles per litre. Also, product *purity* limits the final concentration of undesired side product, *D*, to be no more than 0.018 moles per litre; and process *profitability* demands obtaining at least 0.058 moles of *C* at the end of the batch. Operating conditions are defined here as feed addition rate $0 \leq F \leq 12.0$, [liters min⁻¹] and duration of the semi-batch period $0 \leq t_1 \leq t_{\text{final}} = 120$. Even though more elaborate operating strategies can be used, this simple operating policy is considered here for the sake of clarity and space. Since the optimal solution should necessarily be within a **ROI** such that $\text{Pr}_{\mathbf{u}}[g(\mathbf{x}_f) \geq 0]$ is maximized, there is no point in reducing the overall batch time to less than the available 120 minutes. As a result, only two variables, *F* and *t*₁, are left for defining the operating policy.

For the present case study, it is assumed that for each modelling run the only available measurements are the final concentrations of species *B*, *C* and *D*. Since the short-lived intermediate *I* cannot be observed, a postulated mechanism that was thought to agree well with the data from modelling runs is:



Using this imperfect model and standard regression techniques, a given data set allows estimating $\hat{k}_1$ and $\hat{k}_2$ to account for both the kinetic behaviour and batch variations due to imperfect temperature control⁴¹. It is worth noting that this model can provide only a local approximation to the true process behaviour. Hence, kinetic parameters $\hat{k}_1$ and $\hat{k}_2$ can optionally be re-estimated as the **ROI** is being shrunk in a given iteration of the evolutionary optimization algorithm of Fig. 11.

The initial data set was made up of the 10 policies shown in Fig. 12(a), each of which was tried twice. Note that the initial policies were judiciously chosen in a **ROI** where the policy has some chances of satisfying the endpoint constraints. After ten new policies have been incorporated into the data set, the **ROI** has been drastically reduced as shown in Fig. 12(b). Fig. 13 depicts the probability

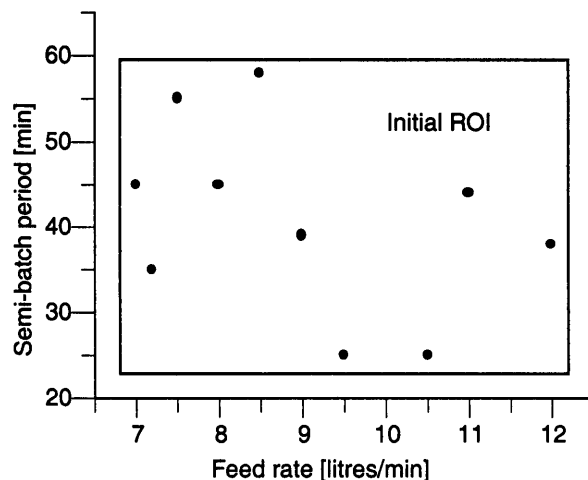


Fig. 12 (a) Initial set of operating policies

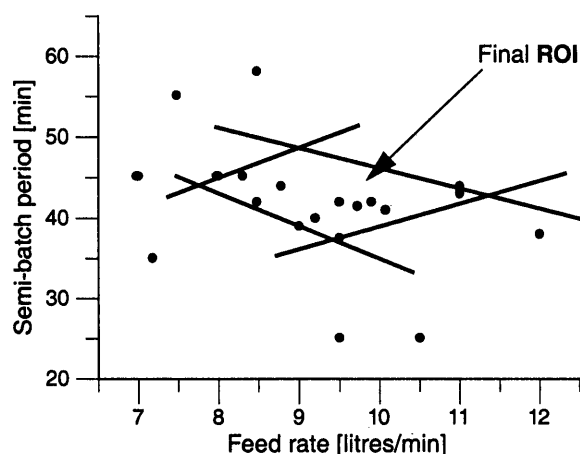


Fig.12 (b) The reduced **ROI** after 10 new policies have been tried

of success estimated using the imperfect model and data collected. The value function for operating policies obtained is shown in Fig. 14. These results were obtained using an exponential “cooling” procedure as follows. After an entirely new control policy for the next run was actually implemented and its resulting information added to DS_{θ} , the temperature parameter τ was decreased geometrically with the cumulative number of control policies *z* explored during learning using:

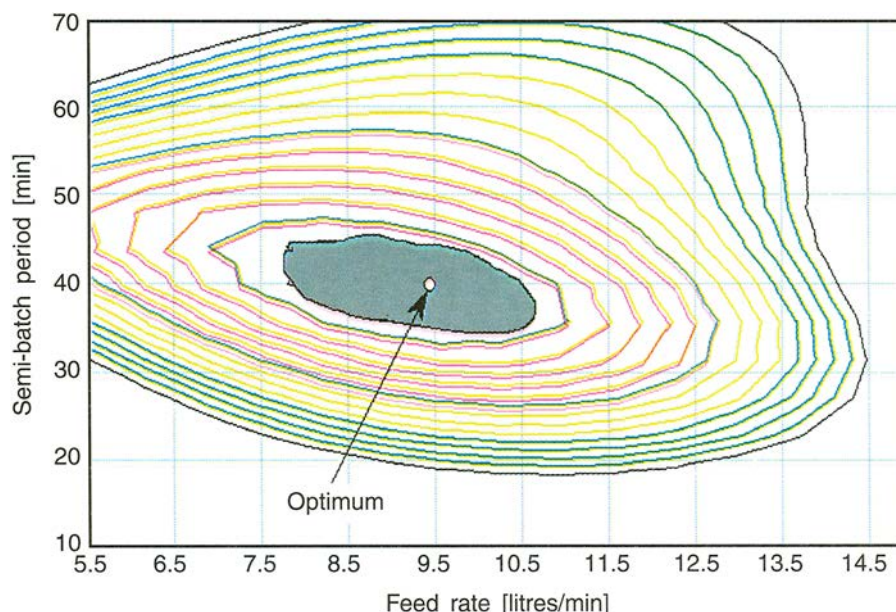


Fig. 13 Probability of success contour plots

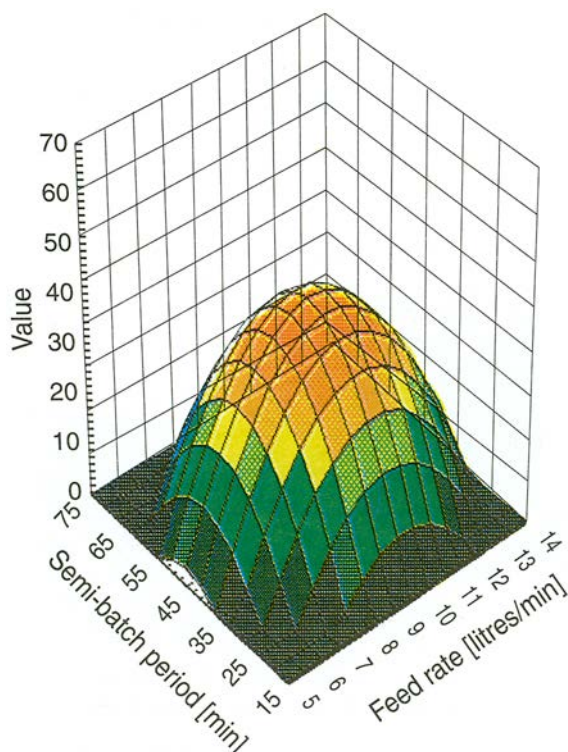


Fig. 14 Value function for the semi-batch reactor in Example 2

$$\tau(z) = p \cdot (1 - p)^z, z \in \{0, 1, 2, \dots\} \quad \dots(25)$$

where $p=0.25$. Note that for $p=0.25$, exploration is still allowed to a certain extent even after 10 new control policies have been incorporated into the set DS_p .

7 Concluding Remarks

Effective experimental design in model development from process data is essential for taking key decisions that can speed up the migration from reactions (in a flask) to shop-floor manufacturing². The successful application of process modelling in evolutionary optimization critically depends on the capability of introducing the required bias when planning each and every experimental run. At the onset of experimentation, when the process is poorly understood, we can only resort to imperfect models and biased data to answer the question: what operating conditions seem worth trying? Ideally, the next experiment should be chosen to obtain the data that minimize the uncertainty about the location of the optimum. Since there exist alternative models for determining an improvement direction, optimization uncertainty is intrinsically related to model discrimination and selection. Considerable savings in money and time can be obtained when the relevant data is obtained at the right time. With the increased emphasis on constant product changeovers and large scale-ups that avoid pilot plant-plant tests, the ubiquitous use of process modelling can make a big difference in the learning curve for process improvement.

In his review, Bonvin⁶ argues that poor measurement and limited process understanding are often more important issues in industry than

any algorithmic aspect referred to the computation of an optimal solution using a model. Moreover, the quest for a feasible and "satisfactory" (i.e. good enough) operating policy often dominates over the notion of absolute optimality, which is a model-related concept. A central concern for any model-based optimization approach for batch units is the necessary alignment between the goal of the modelling effort and the relevance of process data gathered for achieving that goal. Even though this viewpoint is of great importance to industry, it has received only marginal attention in academic research. The notion of modelling for optimization in the face of uncertainty that has been discussed here emphasizes that process models, although imperfect, can still provide useful guidelines for sequentially biasing experimentation. To this end, the importance of active learning from imperfect models and biased/noisy data has been highlighted.

The challenge of modelling for optimization in the process development lifecycle is achieving higher levels of *autonomy*. The availability of high-throughput experimental workstations and

laboratory robotics is shifting the bottleneck to human intervention. The more automated experimental planning for evolutionary optimization would become, the faster process development will be. However, experiment design techniques, e.g. response surface, Taguchi, etc. are not algorithms amenable for automated usage, but rather manufacturing protocols that require human judgment to be involved in the decision making loop. Active learning from imperfect models and noisy data is the approach that has been proposed to introduce automation in experimental optimization, which allows human analytical capabilities to concentrate on more complex and ill-defined issues. More theoretical work is needed to design optimization algorithms that are better equipped to deal with process data rather than numbers. In this regard, the trade-off between exploration and exploitation in modelling for optimization is a key concept that should be better understood. Moreover, industrialists should be more willing to accept that the only way to exploit more is to explore more. There is no an easy way out of the learning curve.

References

- 1 T G Oberg and S N Deming Find Optimum Operating Conditions *Fast Chemical Engineering Progress* April issue (2000) 53
- 2 G Stephanopoulos *et al.* *Comp Chem Engng* **23** (1999) S975
- 3 S Rahman and S Palanki *AIChE J* **42** (1998) 2869
- 4 P Terwiesch, M Agarwal and D W T Rippin *J Proc Control* **4** (1994) 238
- 5 B Schenker and M Agarwal *J Process Control* **5** (1995) 329
- 6 D Bonvin *J Proc Control* **8** (1998) 355
- 7 E C Martinez *Computers and Chem Engng* **24** (2000) 1187
- 8 C Filippi, J Graffe, J Bordet, J Villiermaux, J Barney, P Bonte and C Georgakis *Chem Engng Sci* **41** (1986) 913
- 9 C Filippi, J Bordet, J Villiermaux, S Marchal-Brassey and C Georgakis *Comp Chem Engng* **13** (1989) 35
- 10 J Uhlemann, M Cabassud, M LeLann, E Borredon and G Cassamatta *Chem Engng Sci* **49** (1994) 3169
- 11 J Hamer *Chem Engng Sci* **44** (1989) 2363
- 12 D Bonvin and D Rippin *Chem Engng Sci* **45** (1990) 3417
- 13 O Prinz *Chemometric Methods for Investigation of Chemical Reaction Systems*. Ph D Thesis ETH Zurich Switzerland (1992)
- 14 P Tsobanakis *Tendency Modelling and Optimization of Fed-batch Fermentations* Ph D Thesis Lehigh University (1994)
- 15 C Georgakis *Entropie* **194** (1995) 34
- 16 J Fotopoulos *Quantification of Model Uncertainty and its Impact on the Optimization and State Estimation of Batch Processes via Tendency Modelling*. Ph D Thesis Lehigh University (1996)
- 17 M Amrhein *Reaction and Flow Variants/Invariants for the Analysis of Chemical Reaction Data*. Ph D Thesis ETH Zurich Switzerland (1998)
- 18 M Amrhein, B Srinivasan and D Bovin *Chem Engng Science* **54** (1999) 579
- 19 J Fotopoulos, C Georgakis and H G Stenger Jr *Chem Engng Sci* **49** (1994) 5533
- 20 J Fotopoulos, C Georgakis and H G Stenger Jr *Chem Engng Sci* **51** (1996) 1899
- 21 E C Martínez and C Georgakis *Proc 3rd Mercosur Congress on Process Systems Engineering*, Santa Fe, Argentina **1** (2001) 7
- 22 E C Martinez, R A Pulley and J A Wilson *Chemical Engineering Research and Design* Trans of IChemE Part A September issue (1998) 711
- 23 J A Wilson and E C Martinez *Applications of Neural Networks and Other Learning Technologies in Process Engineering*, (Eds. I M Mujtaba and M A Hussain) Imperial College Press (2001) 269
- 24 J Valappil *Nonlinear Model Predictive Control of End-use Properties in Batch Reactors with Consideration of Uncertainty*. Ph. D. Thesis Lehigh University (2000) (in preparation)

- 25 S Verwater-Lukszo A Practical Approach to Recipe Improvement and Optimization in the Batch Processing Industry Ph D Thesis Eindhoven Technical University The Netherlands (1995)
- 26 S Bijlsma, D J Louwerse and A K Smilde *AIChE J* **44** (1998) 2713
- 27 S B Gadewar, M F Doherty and M F Malone *Comp Chem Engng* **25** (2001) 1199
- 28 K V Waller and P M Makila *Industrial and Engineering Chemistry Process Design and Research* **20** (1981) 1
- 29 R Defay Azéotropisme-Équations Nouvelles de états Indifferents *Bulletin de la Classe des Sciences* **17** (1931) 940
- 30 Y Sedrati *et al.* *Comp Chem Engng* **23** (1999) S427
- 31 E C Martínez and G A Pérez Evolutionary Optimization of Noisy and Costly Functions using Non-parametric Statistical Methods *Proc 3rd Mercosur Congress on Process Systems Engineering* Santa Fe, Argentina **1** (2001) 385
- 32 P Sprent and N C Smeeton *Applied Nonparametric Statistical Methods* Chapman & Hall/CRC London (2000)
- 33 M Hollander and D A Wolfe *Nonparametric Statistical Methods* 2nd Edition John Wiley & Sons New York (1999)
- 34 J V Beck and K J Arnold *Parameter Estimation in Engineering and Science* Wiley New York (1977)
- 35 A C Atkinson and A N Donev *Optimum Experimental Designs* Oxford University Press New York (1992)
- 36 G K Raju and C L Cooney *AIChE J* **44** (1998) 2199
- 37 D Ruppen, D Bonvin and D W T Rippin *Comp Chem Engng* **22** (1998) 185
- 38 P Terwiesch, D Ravemark, B Schenker and D W T Rippin *Comp Chem Engng* **22** (1998) 201
- 39 F P Bernardo, E N Pistikopoulos and P Saraiva *Ind Engng Chem Res* **38** (1999) 3056
- 40 E Walter and L Pronzato *Automatica* **26** (1990) 195
- 41 M Young *Processing* **3** (1980) 27