



PONTIFICIA UNIVERSIDAD CATOLICA DE CHILE
SCHOOL OF ENGINEERING

STRATEGIC REPARAMETRIZATION AS A MODEL ENHANCER: APPLICATION TO WINE FERMENTER DIGITAL TWINS

CRISTÓBAL LUIS TORREALBA VÁSQUEZ

Thesis submitted to the Office of Graduate Studies in partial
fulfillment of the requirements of the Degree of Master's in
Engineering Sciences

Advisor:

JOSÉ RICARDO PÉREZ.

Santiago de Chile, June, 2021

© 2021, Cristóbal Torrealba Vásquez



PONTIFICIA UNIVERSIDAD CATOLICA DE CHILE
ESCUELA DE INGENIERIA

STRATEGIC REPARAMETRIZATION AS A MODEL ENHANCER: APPLICATION TO WINE FERMENTER DIGITAL TWINS

CRISTÓBAL LUIS TORREALBA VÁSQUEZ

Members of the Committee:

J. RICARDO PÉREZ

PEDRO SAA

RICARDO LUNA

LEONARDO VANZI

Thesis submitted to the Office of Research and Graduate Studies in
partial fulfillment of the requirements for the Degree of Master's in
Engineering Sciences

Santiago de Chile, June, 2021

*To my parents, family, and friends,
who with their unconditional love
and support have always enabled me
to climb even the highest of
mountains.*

ACKNOWLEDGEMENTS

In the first place, I would like to express my appreciation and admiration to my supervisor, Prof. José Ricardo Pérez Correa, whom I could always rely on for technical, professional, and personal guidance, and moreover, made possible my insertion at Concha y Toro's Research and Innovation Center to conduct the work described in this thesis. For this, I will be on eternal gratitude.

Also, I would like to acknowledge staff from Concha y Toro's Research and Innovation Center, especially those belonging to the engineering team. Thank you for such a good experience, where an excellent environment, leadership, and camaraderie were the cornerstones of the projects being conducted by this institute. I really appreciate the joyful moments I had with this team, and I will never forget the great moments lived which made part of my first work experience. I would like to give special recognition to Matías Sanz, Bruno Martins, Ana María Betancourt, Carlos Caris, and Mauricio Araya; thank you for your constant positive energies, friendship, and collaborations, which really motivated me during my stay in Talca. Moreover, I would like to acknowledge the engineering team's leader Jose Cuevas, who I deeply admire and respect for his great leadership, career, and friendship; I have learnt some valuable lessons from you, and for that, I will always be in debt to you. In a special manner, I would like to thank Ricardo Luna. Ricardo represents an important person for me, as he inspired me greatly since his work assisting Professor Pérez at the university, motivating me to follow his steps as assistant and researcher. I would like to thank you for the mentorship, dedication, and friendship you gave me during my work at the research center, which I will always bear with me wherever I decide to make my future.

Finally, I would like to express my gratitude to my family, to my parents Javier Torrealba and Macarena Vásquez, who have always loved and supported me since ever. Also, I would like to give special thanks to my love Camila; thank you for your unconditional love and support, for your always necessary corrections of my writing, for your constant motivational speeches, among countless others. You are the most beautiful thing in my life, and for that, I will always be thankful.

LIST OF CONTENT

DEDICATION	ii
ACKNOWLEDGEMENTS	iii
TABLE INDEX.....	vi
FIGURE INDEX.....	vii
ABSTRACT	viii
RESUMEN.....	ix
1. Introduction.....	10
1.1 Motivation	10
1.1.1 Mission of Concha y Toro's Center for Research and Innovation ..	10
1.1.2 CORFO R+D+i project Portfolio	10
1.1.3 Sub-project R+D1	12
1.1.4 Sub-project R+D2	13
1.1.5 Sub-project Prototype.....	14
1.1.6 Sub-project VP	16
1.2 First principles-based models in SmartWinery	17
1.2.1 FPM modeling and simulation	18
1.2.2 FPM modeling applied to winemaking	19
1.3 Kinetic model calibration	27
1.3.1 Dimensions of model robustness.....	28
1.3.2 Strategic reparametrization as a model enhancer	33
1.4 Approach of the thesis.....	36
1.4.1 Hypothesis and objectives.....	38
1.4.2 General methodology	39
1.4.3 Principal results and conclusions	40

2.	Enhancing wine fermentation models in the presence of limited data: a multi-criteria guided parametric robust assessment workflow	42
2.1	Abstract	42
2.2	Introduction	43
2.3	Materials and methods	45
2.3.1	Experimental design	45
2.3.2	Workflow methodology	48
2.3.3	Model validation	59
2.4	Results and discussion.....	60
2.4.1	Initial evaluation of Coleman and Zenteno models	60
2.4.2	Best-overall model structure selection	64
2.4.3	Pilot-scale validation	70
2.5	Conclusion.....	73
2.6	Acknowledgements	74
3.	Final remarks and future perspectives	75
	REFERENCES.....	77
	APPENDIX.....	84
	Appendix A : Reduced Zenteno Model.....	85
	Appendix B : Model structure Zenteno-10325	87
	Appendix C : MCDM results when applied to group I.....	88
	Appendix D : Coleman model structural identifiability analysis results	89

TABLE INDEX

Table 2-1: MCDM methods and evaluated scenarios involved in the viable structure selection stage	55
Table 2-2: Experimental data sets	56
Table 2-3: Predictive performance over <i>laboratory-scale</i> and <i>pilot-scale</i> experimental data when simulating using the original Zenteno and Coleman models	62
Table 2-4: Selected viable structure groups specifications	65
Table 2-5: Model structure Zenteno 10325 averaged performance indices when applied for predictions over the <i>pilot-scale</i> experiments	70
Table B-1: Parameter enumeration and estimation corresponding to the reduced model Zenteno 10325	87

FIGURE INDEX

Figure 1-1: Structure of R+D+i portfolio actually being developed by the QI and IAI teams at CRI.....	12
Figure 1-2: SmartWinery’s global structure, information flow, and modules with their corresponding sub-projects.	14
Figure 1-3: The seven layers of IIOT infrastructure involved with standard application and digital twin design practices	15
Figure 1-4: Graphic representation of interactions within principal components in AF kinetics	23
Figure 1-5: Graphic representation of the reparametrization process	35
Figure 2-1: Laboratory and pilot scale batch experimental setups used for the experiments	47
Figure 2-2: Graphical representation of the workflow used for model robustness assessment and selection.	51
Figure 2-3: YAN corrections applied during simulations.....	63
Figure 2-4: MCDM results when applied to model structures from group II	67
Figure 2-5: Mean absolute values of the components of the last singular vector of the ROSM obtained from applying the method proposed by Stigter & Molenaar (2015) over the reduced Zenteno model	69
Figure 2-6: Validation of model structure Zenteno 10325 applied over the pilot-scale validation experiments.	72
Figure C-1: MCDM results when applied to model structures from group I.....	88
Figure D-1: Mean absolute values of the components of the last singular vector of the ROSM obtained from applying the method proposed by Stigter & Molenaar (2015) over the Coleman model.....	89

ABSTRACT

Dynamic process modeling is a discipline increasing its participation in industrial solutions, nowadays playing part in a several industries for control and prediction purposes in production processes. The correct use of these models depends on parameter estimation; a necessary task for these tools to represent closely a given system. However, enough high-quality data is not always available to carry out this work; scarce sample availability, inadequate sampling times, and high amounts of noise are common conditions impairing dynamic parameter estimation in most calibration procedures. The above may lead into serious difficulties during parameter estimation leading into inappropriate calibrations; if this is the case, highly non-linear models will tend to instability, thus, failing at generating reliable predictions. Therefore, a method to guarantee reliable calibrations when limited data is at hand is necessary if industrial application of these models is sought. In this work, a new method for robustness-guided model reparametrization was developed in collaboration with the Center for Research and Innovation at Viña Concha y Toro to support model calibration. As a case study, alcoholic fermentation models of *Cabernet Sauvignon* wines were recalibrated and improved using the proposed framework. To this task, wine fermentations were conducted in laboratory and pilot scale reactors, monitoring changes in the main metabolites implied in this process by using spectrophotometry. Subsequently, using laboratory-scale data, the developed method was applied to select the model structure that best adapted to the available data, which was then validated using pilot-scale data. By implementing this reparametrization strategy, significant improvements in prediction quality and consistency were achieved, leading into higher prediction fidelity.

Keywords: Batch fermentation, robust control, winemaking, systems biology, predictive model

RESUMEN

El uso de modelos basados en primeros principios es una disciplina que cada día se inserta más dentro de aplicaciones industriales, validando su uso para tareas relacionadas al control y predicciones en procesos productivos. Para el correcto uso de estos modelos, siempre debe de llevarse a cabo una estimación dinámica de parámetros, de modo que estas herramientas representen de manera realista un sistema a caracterizar. Sin embargo, no siempre se tiene una cantidad suficiente de datos de alta calidad para realizar esta labor; la estimación de parámetros de modelos biológicos se caracteriza por una baja disponibilidad de muestras, desconocimiento de tiempos de muestreo adecuados, y altos niveles de ruido en las mediciones. La prevalencia de las anteriores dificultades en datos de diferentes escalas experimentales suele resultar en calibraciones inapropiadas, donde modelos altamente no-lineales e inestables tienden a generar predicciones poco confiables. En esta línea, un método para garantizar la confiabilidad de una calibración se hace necesaria si se busca una implementación industrial de estos modelos. En este trabajo llevado a cabo en el Centro de Investigación y Desarrollo de Viña Concha y Toro, se desarrolló un novedoso método de reparametrización de modelos guiada por indicadores de robustez para apoyar la tarea de la calibración de estos. Este procedimiento fue aplicado al mejoramiento de modelos de fermentación alcohólica de vinos *Cabernet Sauvignon*. Se llevaron a cabo fermentaciones en reactores a escala laboratorio y piloto, donde mediante espectrofotometría se realizó seguimiento sobre los principales metabolitos involucrados en este proceso. Posteriormente, utilizando los datos escala laboratorio, se utilizó el método desarrollado para seleccionar la estructura de modelo que mejor se adaptase los datos disponibles, siendo finalmente validado en la escala piloto. El método permitió obtener modelos más confiables que mejoraron significativamente la precisión y consistencia de sus predicciones.

Palabras Claves: fermentación batch, Control robusto, enología, biosistemas, modelos predictivos

1. INTRODUCTION

1.1 Motivation

1.1.1 Mission of Concha y Toro's Center for Research and Innovation

This thesis is a product of a co-development work, where the collaborators were the Pontificia Universidad Católica de Chile and the Center for Research and Innovation (CRI) at Viña Concha y Toro. CRI is an institution conceived by the changes and new challenges presented by the global wine industry during the last years. CRI's mission is to integrate real industry solutions based on applied investigation, innovation, and technology transfer to promote Concha y Toro's productive excellence and sustainability, aiming to impact the national and international wine industry. According to CRI's Strategic Plan for Research and Development (2016-2020), five research areas have been defined given their importance over the productive chain of the company. These areas are: i) Strengthening of the vegetal production area; ii) Management of hydric and scarce resources; iii) Quality index for grapes and wines (QI); iv) Instrumentation, automation, and insertion of TI (IAI); v) New products design. Each area has its own associated development team and associated projects, where challenges are approached interdisciplinarily.

1.1.2 CORFO R+D+i project portfolio

Regarding the objective of this thesis, teams from strategic areas QI and IAI have been designated to lead a series of projects that are part of an R+D+i project portfolio sponsored by CORFO and Concha y Toro since 2017. This portfolio consists in a total of four sub-projects (Figure 1-1): Prototype, R+D1, R+D2, and Validation and Packaging (VP). Though presenting different specific objectives, all the previous sub-projects aim to a sole final purpose, which is the creation of a

digital platform for the management of grape and wine quality named SmartWinery. The SmartWinery platform intends to present itself to oenologists as an instrument for the maximization of efficiency during wine processing decision-making, ultimately looking for an overall increase in profit per liter of wine produced. The last also includes aiding wine professionals in the adaptation to seasonal changes affecting winemaking (such as climatic change), which represents one of the major challenges in this industry nowadays (Fraga, 2020). To accomplish its main purpose, the SmartWinery platform approaches the multifactorial problem that characterizes the production of wine focusing in optimizing management and resource use during processing. Specifically, this production optimization involves looking into each link in the company's productive value chain, supporting decision-making from its initial stages, where fruit is received and classified, until the final stage, where wine is bottled. Here, decrease of the frequency of wine reclassification (i.e., wine originated from high-quality grapes being downgraded in quality terms given an inefficient processing), use of fermenters optimization for productivity increase while reducing operational costs, upholding wine quality consistency, among others, are some of the objectives supporting the SmartWinery's development. All these objectives are approached in an interdisciplinary way, where knowledge fields such as process engineering, data science, agronomy, biotechnology, among many others, meet to guide the development of this new technology.

To ensure that these objectives are fulfilled, two main modules make up the SmartWinery, which are addressed in sub-projects R+D1 and R+D2, while sub-project Prototype involves mounting the previous modules in a digital infrastructure to generate a Minimum Viable Product (MVP). Posteriorly, this MVP is furtherly rectified and enhanced as part of sub-project VP. A brief explanation for each subproject is given in the following subsections.

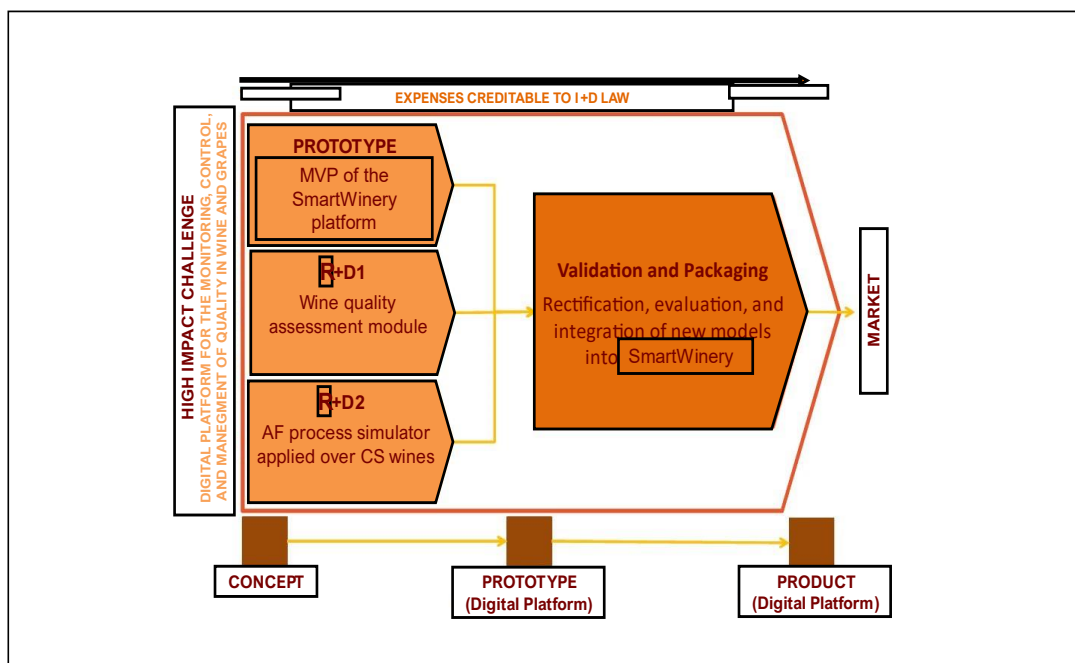


Figure 1-1. Structure of R+D+i portfolio being developed by the QI and IAI teams at CRI.

1.1.3 Sub-project R+D1

Regarding the sub-project R+D1, its main objective is the creation of a module for wine quality evaluation; a complex task given the difficulty to define what is “good quality” in quantitative terms during winemaking, as there is no clear relationship between many of wine’s attributes and its acceptance (Brossard et al., 2016; Dinnella et al., 2011). To do so, a combination involving wine characterization using advanced laboratory techniques to detect chemical markers (mainly spectrophotometry), expert-panel guided wine tasting and evaluation, and application of sophisticated data analysis and metanalysis techniques are joined together in the quest of defining which chemical markers are correlated with the concept of “high-quality” in winemaking. If the above is accomplished, then there would be a rational way to qualify the wines produced in Concha y Toro’s wineries, thus, implying that an optimal composition of this beverage is possible to

define, ultimately fulfilling the purpose of guiding wine production towards maximum-quality.

1.1.4 Sub-project R+D2

Subproject R+D2 aims to the creation of a wine processing simulator. Specifically, this is being focused on characterizing and predicting changes occurring in wine fermenters during the Alcoholic Fermentation (AF) process, showing how the main metabolites of the fermenting wine change throughout the process in function of the external inputs associated on how reactors are being operated. To do so, properly calibrated first principles-based modeling (FPM) is employed by applying sets of differential and algebraic equations (DAE) representing mass and energy balances in the system. In this context, many AF models are available to accomplish the previous task; however, these represent generic approaches to the vinification process, existing many factors involved in the outcome of this process (e.g., grape composition, yeast genetics, among others), which are typically unique for each productor. The above involves that, whichever model is selected for process simulation, a proper calibration is imperative, as this procedure will be able to capture the effect of factors affecting a vinification procedure, consequently guaranteeing its adequate representation by the model. In the context of the SmartWinery platform, well-mixed wine AF models (discussed in section 1.2) have been calibrated and implemented in the platform using data from laboratory-scale fermentations. However (and serving as purpose to this thesis), there is still space to enhance these models, which still require a further validation in larger process scales. Sections 1.2 and 1.3 give greater detail about the concepts inquired to address the last situations, while chapter 2 gives a detailed summary about work performed to accomplish them.

1.1.5 Sub-project Prototype

Subproject Prototype focuses on the creation of the digital platform serving as MVP of SmartWinery, designed to implement the modules generated from sub-projects R+D1 and R+D2. Here, collaboration with the different vineyards and wineries belonging to the Concha y Toro company, and incorporation of IIOT technologies, aim to predict the behavior of industrial fermentations throughout processing, as well as concentration of tannins and anthocyanins in finished wines; widely used as quality assessment indices in this industry (Holt et al., 2008; Ma et al., 2014; Singleton & Trousdale, 1992). The modules by themselves accomplish a specific task at simulating, predicting, and characterizing the wine production process; however, their true potential relies on what occurs when they are suitably combined. As shown in Figure 1-2, it is possible to communicate each of these modules so that the information flow is established systematically between each of the SmartWinery's functionalities.

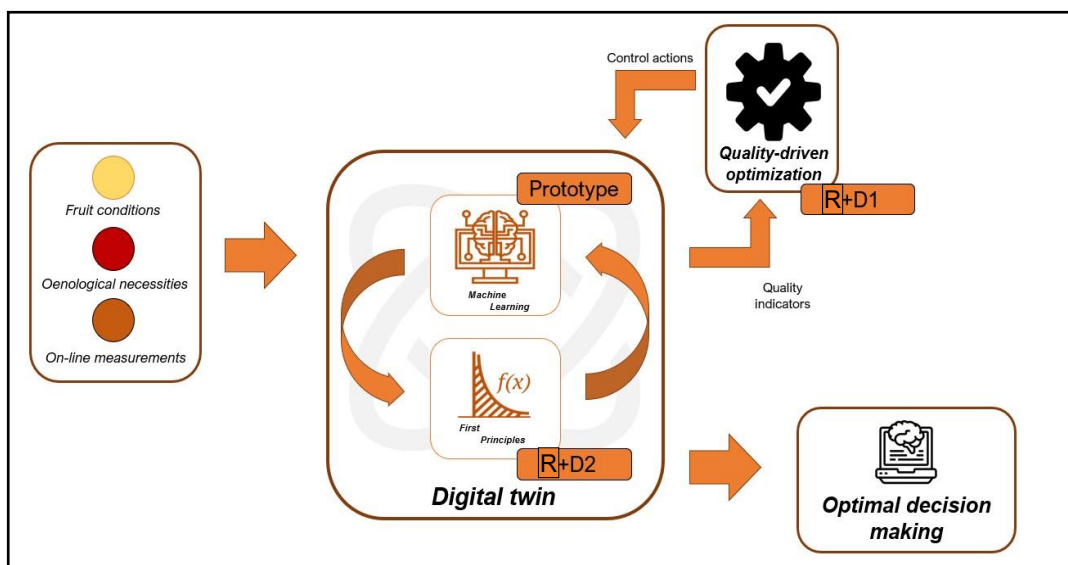


Figure 1-2. SmartWinery's global structure, information flow, and modules with their corresponding sub-projects.

A configuration with these characteristics, supported by an adequate digital infrastructure (Figure 1-3), are the fundamental building blocks to generate a digital twin, i.e., a virtual representation of an object or system that spans its lifecycle, is updated from real-time data, and uses simulation, machine learning and reasoning to help decision making (Armstrong, 2020). Given their enabling features and great versatility, digital twins represent a trailblazing technology, promising extreme usefulness if smartly used. This has already been proven in other industries, especially in those aiming towards smart manufacturing (Tao et al., 2019). Considering the above, a major impact over the winemaking industry is expected with the creation and implementation of digital twins of wine fermenters, which, by using modelling and live measurements, can be employed for optimization in a new and further extent; thus, enabling to accomplish the previously stated objectives supporting the SmartWinery's development.

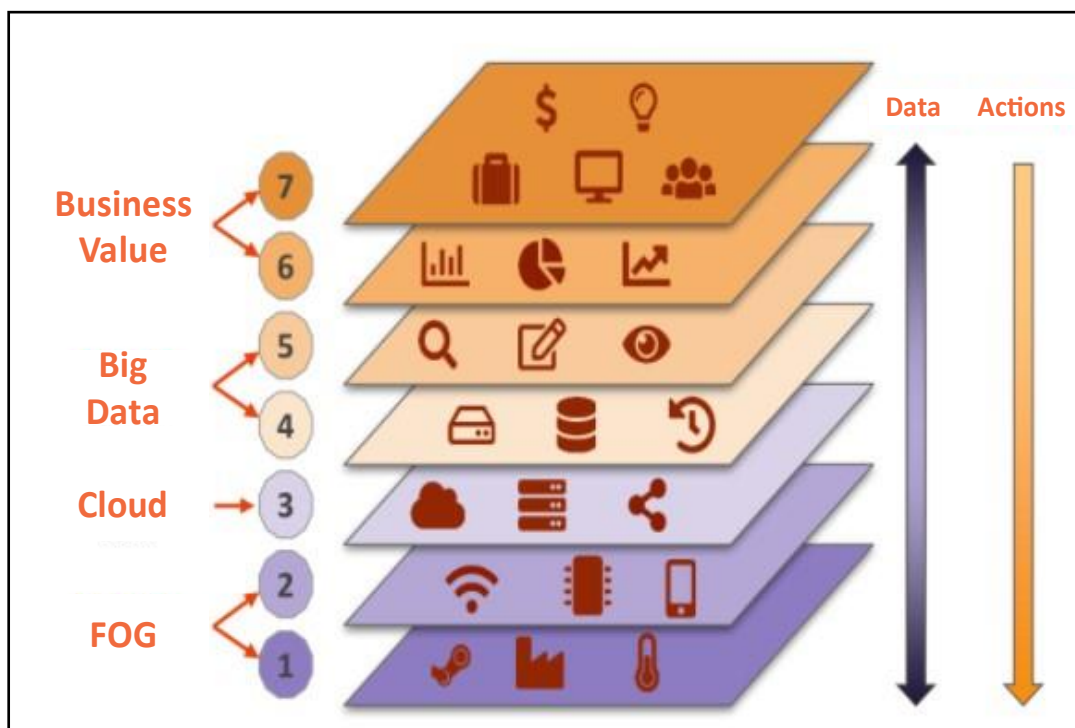


Figure 1-3. The seven layers of IIOT infrastructure involved with standard application and digital twin design practices. Taking SmartWinery as an example,

fog-level layers include sensors, controllers, and actuators installed in fermentation tanks. The *Cloud* layer refers to those technologies used for storage of big amounts of data (servers), while *Big-Data* layers are those associated to the different computational tools used for modeling and prediction generation. Finally, the *Business-value* layers include technologies for automatically consuming/processing information of the previous layers, overall aiming to aid in decision-making in wine fermentation management. Adapted from Luna et al., (2020).

1.1.6 Sub-project VP

Sub-project VP sustains from the other sub-projects, as its main objective is the evaluation, integration, and further application of the SmartWinery platform in industrial fermentations for its validation. The above includes the improvement over the re-design and extension of the TI structure of SmartWinery, implementation of the digital platform into more wineries in the company, enhancement and extension of functionalities involving models developed in subprojects R+D1 and R+D2, techno-economical evaluation of the modules in the SmartWinery, among others. All the activities mentioned above were performed mainly in direct collaboration with staff working at Lourdes winery from Concha y Toro, where the implementation of IIOT hardware, as well as SmartWinery's software were conducted to make the above possible. Here, several industrial-scale wine fermenters were adapted to a series of digital technologies (web applications, sensors, among others) and coupled to the SmartWinery, to finally be managed and evaluated by oenologists from this establishment. Also, the CRI experimental winery was included in this assessment, where the same methodology was implemented in pilot-scale fermenters for performance evaluation of the SmartWinery.

As the readers may infer from previous subsections, the development of the SmartWinery platform is a project of major proportions, thus, requiring the support of several developers for its construction. Following this idea, this work positions itself in sub-project VP, having its main objective related to the task of further calibrating and enhancing models associated with this digital platform, specifically aiming to models based on first principles (those developed in sub-project R+D2). In the next sections of this chapter, it will be discussed how these models fit in the construction of the SmartWinery, as well as how their functionalities are oriented into collaborating with the other type of models. Posteriorly, fundamental concepts about FPM and parametric robustness will be discussed, as they play a key role over the predictive capacity in this type of model, thus, are essential if there is an intention of using these models consistently in the SmartWinery. Finally, the approach of this thesis, including hypothesis, methodology, results, and conclusions will be stated to give readers a solid idea of the main objectives behind the work in Chapter 2, which presents an article to be published in an international journal describing in detail this thesis's work.

1.2 First principles-based models in SmartWinery

To understand how FPM models fit in the development of the SmartWinery, we must first understand their origins and how they work. Wine fermentation is one of the oldest bioprocesses known by the humankind, thus, the knowledge regarding this process has grown exponentially since its discovery. The above extends especially into the last century, where, given the uprise of knowledge in fields related to chemical process engineering (e.g., chemical kinetics and heat transfer), and furtherly complemented with the emergence of enabling technologies for process simulation, the understanding of the fermentation phenomena was revolutionized (Miller & Block, 2020). This naturally impacted the comprehension in winemaking-related research, as the AF process is the heart of any winery; most of the extraction and bioconversion phenomena occur at this stage of the process

(Unterkofer et al., 2020). The high impact of AF over the final product's quality during winemaking has driven the interest of authors into channeling this new knowledge into the development of instrumentation valuable for wine production optimization. In this context, mathematical modeling has risen as a powerful tool in this industry, as it shifts the paradigm of practical trial-and-error into a much more cost-effective digital approach, ultimately enabling the insertion of technologies originated from other disciplines in engineering (such as process control and plant-scale optimization). A description referencing to mathematical modeling of the AF process will be presented in the following subsections for a deeper understanding regarding to these tools.

1.2.1 FPM modeling and simulation

FPM has proved to be an outstanding tool in the field of chemical engineering, especially given its reliability, flexibility, and relatively simple structures when applied to most experimental systems. Its range of applications is wide: real-time optimization; model predictive control; process performance monitoring; closed-loop control and automation; among others (Pantelides & Renfro, 2013). These last functionalities have proven great potential across the process industry, where we find their application in different productive areas such as the food and pharmaceutical industries (Benyahia et al., 2012; Mahdi et al., 2009). These models are constructed using fundamental engineering, physics, and chemistry principles, thus, differing from those derived purely from plant and/or other data (e.g., time-series, neural networks, and other forms of data-based models). Generally, these fundamental principles rely on mass and energy balances, which employ physical/chemical originated terms to include the effect of phenomena associated with concepts such as thermodynamics, chemical kinetics, among many others (Pantelides & Renfro, 2013). Given their nature, FPM models have a dynamic behavior, meaning that the processes they model are constructed

transiently and continuously throughout time; another advantage when compared to purely data regression-based models.

A typical FPM model is described by a dynamic system of differential equations (ODEs), as presented in Equation 1.1 and Equation 1.2. Here, f denotes the dynamic model structure, h the output observation function, $x(t)$ the vectorized model states, $u(t)$ the vectorized inputs, $y(t)$ the vector of measured/observed outputs, and θ the vector of model parameters.

$$\frac{dx(t)}{dt} = f(x(t), u(t), \theta) \quad (1.1)$$

$$y(t) = h(x(t), \theta) \quad (1.2)$$

The system above, while in some cases too simple (e.g., large systems presenting significant concentration/temperature gradients), may be applied to most dynamic systems for generating predictions while considering external inputs and changes occurring throughout processing. In mathematical terms, we refer to simulation as the integration of the ODE system presented in equations 1.1 and 1.2. However, the outcome of this operation depends highly on how the system is perturbed (effect of external input $u(t)$, as well as the values assigned to parameters θ), as these characterize fundamental processes occurring in the system in the form of constants that represent the effect of conversion rates, heat transfer coefficients, activation energy, among others.

1.2.2 FPM modeling applied to winemaking

Approaching the winemaking process, the system described in equations 1.1 and 1.2 must be adapted so that the nature of the fermentation process is adequately represented. Fermentation is a process where microorganisms and their enzymes bring about desirable changes over a specific matrix (typically a food matrix), where one or more substrates are bioconverted into a wide range of possible

metabolites depending on the characteristics of the biological system employed. In wines, grape must represent a complex media composed of a rich blend of amino acids, sugars, organic acids, and so on. This richness represents an ideal culture media for most microorganisms, thus, encouraging winemakers to select specific biological systems to ferment grape must in a controlled and reproducible way, avoiding the growth of non-desirable microorganisms (Henriques et al., 2018). The main organism used for AF during winemaking is the yeast *Saccharomyces cerevisiae*, as it has shown superior aptitudes when metabolizing media with a high amount of sugar, and relatively low concentrations of nitrogenous compounds, as well as presenting significant resistance to temperature changes and sulfur dioxide (Suárez-Lepe & Morata, 2012). The most significant advance in the application of FPM on fermentation systems is attributed to the “Monod” kinetic model of cell growth, which despite being proposed several decades ago, still takes part in most of the state-of-the-art growth models (Dette et al., 2005; Miller & Block, 2020; Monod, 1949). The Monod model considers substrate S as a regulator of cell growth acting as a function of total biomass concentration X and a specific growth rate μ (Eq. 1.3). Equations 1.4 and 1.5 incorporate additional terms (K_S and $Y_{X/S}$) to show how the Monod model approaches substrate consumption and biomass specific growth rate, respectively.

$$\frac{dX}{dt} = \mu X \quad (1.3)$$

$$\frac{dS}{dt} = -\frac{\mu X}{Y_{X/S}} \quad (1.4)$$

$$\mu = \frac{\mu_{max} S}{K_S + S} \quad (1.5)$$

Though the Monod model can be used for wine fermentation modeling into some extent, this model presents several limitations in this task given its simplicity: no product generation rate equations; lacking secondary phenomena such as inhibition by product concentration; non-consideration of operational variables such as temperature; single substrate consumption, amongst others (Miller & Block,

2020). Lack of inclusion of the above implies a limitation in the predictive capacity of the model, thus, leading to lower reliability and consistency when performing simulations and comparing with empirical wine fermentations. A more realistic representation of yeast metabolism occurring during the AF in winemaking is shown in Figure 1-4. Here, the fluxes orientation of the principal metabolites involved in the AF process are represented with arrows, while inhibitory effects are shown as T-end lines. We identify glucose (Glx), fructose (F), and nitrogen (N) as the three main substrates in this type of fermentation, which are consumed for the growth and maintenance of yeast cells. On the other hand, the outputs generated in AF by yeasts are mainly ethanol (E), biomass (X), carbon dioxide (CO₂), glycerol (G), and acetic acid (Ac). As shown in the figure, and supported by several researchers, the increase in the concentration of ethanol in fermentation media shows an inhibitory effect over sugar consumption, leading to the death of yeasts given their incapability of obtaining nutrients for maintenance (Brown et al., 1981; Holzberg et al., 1967; Zhang et al., 2015). Though the Monod model approach is much simpler and thus easier to apply, the multiplicity of interactions present in a multi-component system (such as wine must) imposes the necessity of more sophisticated models, where inter-component interactions should be represented by incorporating new expressions that account for them. In this topic, several authors have proposed more adequate model structures, showing a significant improvement when compared to the Monod model. As suggested by Miller & Block (2020), three conceptual categories make up the scope of models approaching the wine fermentation process:

a) Models assuming well-mixed fermentation kinetics

These models essentially consist of systems of DAEs as those shown in equations 1.1 and 1.2. Here, models assume fermentation dynamics work homogeneously in all the volume of fermenters, neglecting any spatial differences. Models in this topic differ by the way they incorporate the effect of operational conditions over changes in the concentration of primary and secondary metabolites, focusing on

different mechanisms related to yeast metabolism to characterize the above. Incorporation of nitrogen as the main limiting nutrient, cellular death as the effect of temperature and ethanol concentration increase, extension of the effect of temperature over kinetic parameters, and incorporation of the effect of oxygen concentration during the initial stage of fermentation are some of the approaches authors have proposed in the mission of modeling well-mixed wine fermentation systems (Boulton, 1980; Cerda-Drago et al., 2016; Coleman et al., 2007; Cramer et al., 2002; Saa et al., 2012). A secondary approach to well-mixed systems relies on the application of metabolic engineering models, where metabolic pathways take a central role in their construction (Pizarro et al., 2007; Sainz et al., 2003; Vargas et al., 2011). Here, a “grey-box” modeling approach is employed by using yeast genetics to define internal fluxes of nutrients which accomplish a specific function inside the metabolic network. Examples of the previous are metabolite importation/exportation, maintenance reactions, DNA synthesis, etc.

Alcoholic Fermentation kinetic models

¿How do they work?

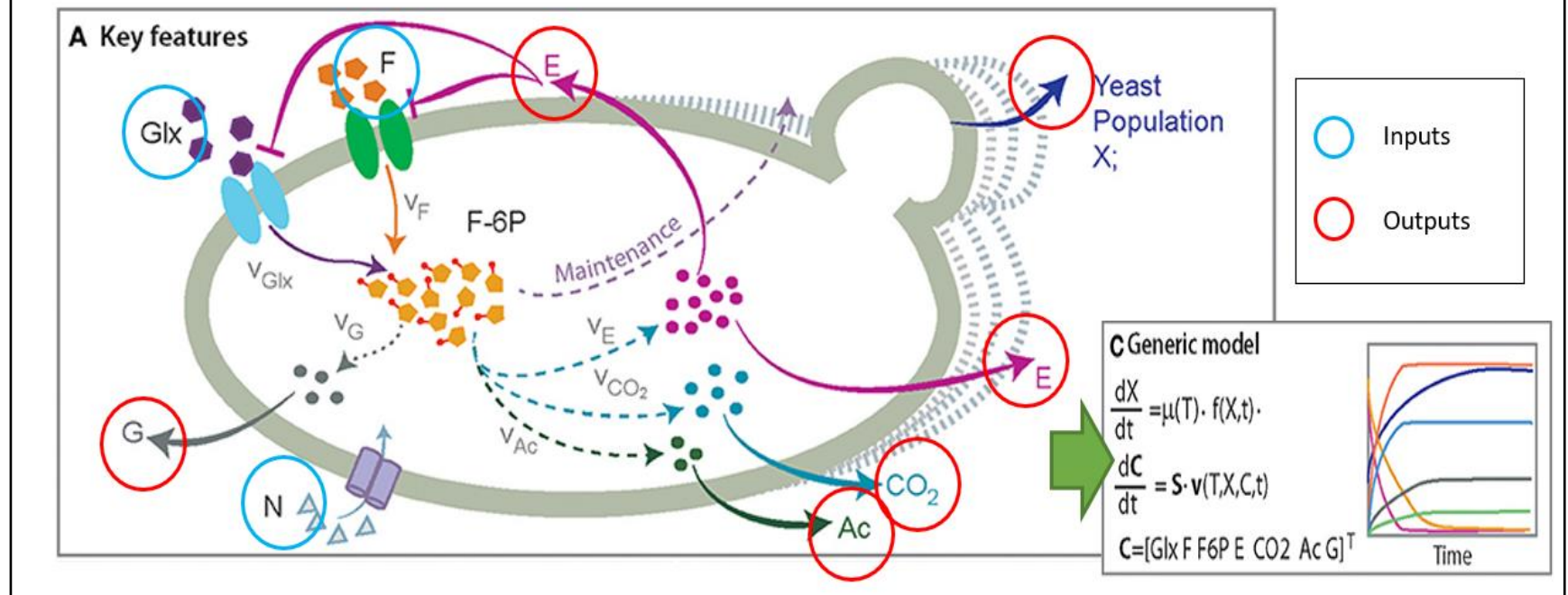


Figure 1-4. Graphic representation of interactions with principal components in AF kinetics. Adapted from Henriques et al., 2018.

Dynamic flux balance analysis (dFBA) takes a central role in metabolic engineering models, as it can determine how yeasts will internally direct metabolite fluxes in time-variant systems, so that a biological objective (biomass production maximization) is accomplished. In this line, several authors have shown interesting results when applied to wine fermentation modeling, showing high precision in the prediction of the generation of secondary metabolites such as glycerol, an important metabolite that affects significantly sensorial attributes (Pizarro et al., 2007; Sainz et al., 2003; Vargas et al., 2011; Sánchez et al., 2014a). Moreover, the latest advances in genomic-scale modeling of *Saccharomyces cerevisiae* yeast species enable exploiting new knowledge on yeast genetic mechanisms to enhance wine secondary metabolite generation, which could lead to a better understanding and further optimization of the wine production process (Sánchez et al., 2017).

b) Models assuming heterogeneous systems

Contrarily to the previous models, this approach does not accept the assumption of the system being well-mixed, which is generally inappropriate for large reactors, where temperature and concentration gradients are significant; this especially applies to industrial winemaking, where these gradients are produced by the presence of solids in red wines (Miller & Block, 2020; Schmid et al., 2009). White wines also exhibit these gradients, as agitation is avoided to control unwanted oxygenations of the must that could lead to detrimental off-flavors (Unterkofer et al., 2020). To better illustrate the generation of gradients during winemaking, different authors have experimentally tested these gradients empirically. Vlassides & Block, (2000) determined that though initially heterogeneous, 1200 L white wine fermenters turned into well-mixed systems rapidly as fermentation rates increased. However, Schwinn et al., (2019) demonstrated that in larger white wine fermenters (7000 L), gradients lasted up to 4 days before becoming well-mixed. This situation is more critical for red wines, where given the generation of a cap (grape pomace agglomeration at must's

surface) temperature could vary up to 10 °C at the solid-liquid interface between the cap and the must (Schmid et al., 2009). The above leads to the conclusion that when large fermenters are involved, significant gradients are expected to develop and persist throughout the process. While displaying a good predictive capacity on smaller systems, all previous models assume well-mixed kinetics, thus non-applicable to large systems. The difficulty that gradient generation represents during winemaking has encouraged authors to create wine fermentation models using more complex mathematic background. An initial effort for the above was proposed by Zenteno et al., (2010), where a compartmentalization method was applied to capture temperature and concentration gradients. More recently, Miller et al., (2019), proposed the use of a computational fluid dynamics software to generate a model with a higher precision (by using smaller compartments) that could better explain the temperature and concentration gradients, successfully predicting biomass and ethanol concentration gradients, as well as describing fluid flow patterns. Developing these models can be challenging; however, they stand-up as the future in AF modeling.

c) Phenolic extraction models

These models mainly apply to red wine fermentations, as only these include a solid-liquid extraction while simultaneously fermenting the sugars in grape must. This extraction is critical during the AF process of red wines, as extracted phenolic compounds are responsible for most of red wine's desirable sensorial properties (Brossard et al., 2016; Unterkofler et al., 2020; Setford et al., 2017). Here, readers must understand that though phenolic extraction does not significantly affect AF kinetics, the contrary is not true (Setford et al., 2017; Setford et al., 2019). Also, both AF and phenolic extraction dynamics share sensitivity to some of the operational variables manipulated during winemaking. For example, in winemaking it has been proven that phenolic compound extraction depends on temperature and solvent properties, the last being dynamic given the constant generation of ethanol because of AF, which is also affected by temperature

(Unterkofler et al., 2020; Yacco et al., 2016). Considering the above, it can be concluded that the AF and phenolic extraction processes involved in winemaking are deeply intertwined, and thus, should be brought together during modeling to be effectively exploited in real-time control systems driven by sensorial quality potential (Setford et al., 2017). This also implies that if large systems are being analyzed, phenolic extraction is bound to be spatial and temporal dependent, as the way fermentation carries out at different locations in fermenters dictates how extraction kinetics will unfold.

Phenolic extraction kinetic models have long been studied differing in the mass transfer mechanisms they consider, as well as the assumptions and fermentation kinetics they employ. While interesting, phenolic extraction models and mechanisms are out of the scope of this thesis, as these depend significantly on the fermentation kinetics, which were the focus of this work. However, for further understanding of phenolic extraction models, it is strongly suggested to consult the work published by Block & Miller (2020). Also, it is suggested to review the work published by Miller et al. (2020), as these authors show results when applying a state-of-the-art spatial gradient-based model coupled with fermentation kinetics obtaining good predictions.

Concluding this section, it has been proven that nowadays there is a good amount of AF process models for winemaking. Actual challenges mainly consist in solving difficulties related to spatial gradients and coupling extraction kinetics in industrial applications, which as addressed previously, has been initially approached in a successful way. Moreover, new challenges stand upon the modeling of the winemaking process, especially the way simulation results are used for value generation when applied with industry representatives. Here, the main difficulty is the lack of knowledge regarding the usefulness and applications where these models can benefit its final users, as most traditional oenologists do not necessarily share a formation oriented to the understanding of this type of complex engineering tools. Moreover, existing models may be robust for simulating wine's

primary metabolites, however, the presented models still have further potential for development, where interesting approaches such as the addition of aroma-related precursors to available state-of-the-art models represents an interesting field to be exploited in the quest of further understanding the winemaking process (Bartsch et al., 2019).

1.3 Kinetic model calibration

In the previous section, AF and phenolic extraction modeling were addressed, determining the existence of several alternatives for the simulation of these processes, each with its unique benefits and limitations. No matter which model is selected, there must always be a calibration step, as FPM models are subject to the user's unique experimental conditions. This means that, in an efficient model, mechanisms affecting each component are correctly explained, however, the magnitude of their effect in the system is variable subject to external conditions. For example, in winemaking, yeast genomics, grape must composition, and nitrogen source are some of the most typical conditions that vary between winemakers. This model calibration step is known as a dynamic parameter estimation (DPE) or “regression” process, where the parameters (θ) of a given model structure (Eq. 1) are calculated using optimization procedures to determine which values fit better the experimental data associated to the process described by the model. However, because experimental data commonly is noisy and incomplete, diagnostics to test model identifiability and validity, and the significance and determinability of their parameters are imperative to determine the degree of experimental support of the model (Jaqaman & Danuser, 2006; Krausch et al., 2019; Saa & Nielsen, 2017). If the calibrated model shows substantial support, then we call the resulting calibrated model a “robust” model structure, which will guarantee reliability while predicting a system's response. In the following subsections, model calibration fundamental concepts will be

described, as calibration and enhancement of the AF models used in the SmartWinery was the main objective guiding the work of this thesis.

1.3.1 Dimensions of model robustness

Robustness is an important condition to address during the DPE process, as it is determinant for the reliable application of models for tasks such as predictive control. The nature of models used for fermentation modeling (i.e., based on first principles and constructed around the concepts of energy and mass balance) imply that these are constructed using a set of rules based on *a priori* hypothesis; thus, regression using experimental data is the natural way for determining their unknown parameters (Jaqaman & Danuser, 2006). However, the achievement of the robustness condition (i.e., a model where all parameters are characterized as significant and determinable, leading into highly reliable predictions given any experimental conditions) during DPE in these models can be challenging. Commonly, for biological system models (as those used for AF) data is scarce and noisy, while models are typically characterized with a high amount of non-linearities (Bonate, 2011; Krausch et al., 2019; Saa & Nielsen, 2017; Sacher et al., 2011). These common data-imparities require a robust model structure so that regression results are accurate, making posterior conclusions reliable (Jaqaman & Danuser, 2006). For robustness assessment, three main stages are approached in this thesis:

a) *A priori* regression diagnostics

The assessment of structural identifiability is a task that must be performed for a given model on an *a priori* basis (prior to regression). This process seeks to answer if, given a specific model structure ($f(x(t), u(t), \theta)$ in Eq. 1.1), all parameters of the model are possible to be determined uniquely with the measured model states (this does not account sampling times nor noise in data). In this task, the main approach is testing the model's output response to changes in the values of its parameters (θ). This is commonly known as a parametric sensitivity analysis, where

parametric output sensitivities can be easily derived for a continuous state-space model with n modeled states and p parameters using the system presented in equations 1.6 and 1.7 (Stigter & Molenaar, 2015).

$$\frac{dx_\theta(t)}{dt} = \frac{\partial f}{\partial x} \cdot x_\theta(t) + \frac{\partial f}{\partial \theta} \quad (1.6)$$

$$y_\theta(t) = \frac{\partial h}{\partial x} \cdot x_\theta(t) + \frac{\partial h}{\partial \theta} \quad (1.7)$$

Here $x_\theta(t)$ denotes sensitivities as a time-dependent ($n \times p$) matrix, with each column containing the sensitivity of each of the model's states to an individual parameter of the θ parameter vector. Similarly, $y_\theta(t)$ represents the matrix of output sensitivities, where only measured states (those in $h(x, t)$ from Eq. 1.2) are included. Equations 1.6 and 1.7 form a linear time-varying system (with Eq. 1.1 and Eq. 1.2), which can be solved simultaneously to obtain output parametric sensitivities (y_θ). Calculated sensitivities are posteriorly used in the construction of the Relative Output Sensitivity Matrix (ROSM), which plays a key-role when analyzing output sensitivities of a system with different physical dimensions. Considering N observations in a time interval $[t_0, t_N]$, the ROSM is then constructed using the definition stated in Eq. 1.8 (Stigter et al., 2015).

$$ROSM(t_N, \theta) = \begin{bmatrix} \frac{\theta_1}{y_1(t_0)} \cdot \frac{\partial y_1(t_0)}{\partial \theta_1} & \dots & \frac{\theta_p}{y_1(t_0)} \cdot \frac{\partial y_1(t_0)}{\partial \theta_p} \\ \vdots & \dots & \vdots \\ \frac{\theta_1}{y_n(t_0)} \cdot \frac{\partial y_n(t_0)}{\partial \theta_1} & \dots & \frac{\theta_p}{y_n(t_0)} \cdot \frac{\partial y_n(t_0)}{\partial \theta_p} \\ \vdots & \dots & \vdots \\ \frac{\theta_1}{y_1(t_N)} \cdot \frac{\partial y_1(t_N)}{\partial \theta_1} & \dots & \frac{\theta_p}{y_1(t_N)} \cdot \frac{\partial y_1(t_N)}{\partial \theta_p} \\ \vdots & \dots & \vdots \\ \frac{\theta_1}{y_n(t_N)} \cdot \frac{\partial y_n(t_N)}{\partial \theta_1} & \dots & \frac{\theta_p}{y_n(t_N)} \cdot \frac{\partial y_n(t_N)}{\partial \theta_p} \end{bmatrix} \quad (1.8)$$

Once the ROSM is constructed, structural identifiability can be assessed. Here, two conditions dictate if a model is structurally identifiable. Firstly, all columns of this matrix must have at least one large entry, reflecting that at least one model

state is being strongly influenced by a given parameter (Jaqaman & Danuser, 2006). Secondly, the ROSM must have full rank; this means columns must be linearly independent, reflecting there is no correlation between parameters. This implies that there is no set of parameters compensating one another when their values change (Jaqaman & Danuser, 2006; Miao et al., 2011; Stigter & Molenaar, 2015). If the conditions above are to be accomplished, then, whichever values are assigned to parameter vector θ , the model will be structurally identifiable.

b) Regression scheme

Once a model is determined to be structurally identifiable (if not, treated so that this condition is achieved, concisely explained in the next subsection), the next step for its application is the regression to determine an optimal set of values for parameters in θ , which minimizes differences between simulated and observed data. To fulfill this, typically maximum likelihood (ML) and least squares (LS) methods are used to approach the regression (Jaqaman & Danuser, 2006). Assuming we aim to determine $\hat{\theta}$, equivalent to the parameter vector which maximizes the probability of observing data y_{obs} when using a determined model ($p(\theta|y_{obs})$), a common statistical model to define the optimization problem when considering independent additive Gaussian noise with constant variance for each measurement is formulated as shown in Eq. 1.9, which is the typical approach used in ML parameter estimation (Saa & Nielsen, 2017).

$$p(\theta|y_{obs}) = (2\pi)^{-N/2} \cdot \det(\Sigma_{meas})^{-\frac{1}{2}} \cdot \exp\left\{-\frac{1}{2}(y_{sim} - y_{obs})\Sigma_{meas}^{-1}(y_{sim} - y_{obs})\right\} \quad (1.9)$$

In the ML approach, the most likely values of θ are those which maximize the probability in the statistical model presented in Eq. 1.9. On the other hand, the LS paradigm addresses the problem stated in Eq. 1.9 by applying the assumption that the error covariance matrix (Σ_{meas}) is constant and independent of θ . Hence, by applying the monotonical log-transformation of Eq. 1.9 under this assumption, this

expression takes the form of the commonly used Residual Sum of Squares (RSS) presented in Eq. 1.10; this function is to be minimized to determine $\hat{\theta}$.

$$RSS = (y_{sim} - y_{obs})^T \Sigma_{meas}^{-1} (y_{sim} - y_{obs}) \quad (1.10)$$

Moreover, other significantly different approaches exist to address the regression problem, such as paradigms based on Bayesian statistics and Monte-Carlo sampling. In the former, parameters are treated as truly random variables, thus, being defined by probabilistic density functions. On the other hand, the above is a sampling-based procedure, which scans the likelihood surface generated from mass simulation for parameter inference. Though interesting, these methods are out of the scope of this work; readers are referred to Saa & Nielsen (2017) for a better comprehension over these methods.

Whichever approach is used, when applied to models related to biological systems, it is almost certain that a significant amount of non-linear relations will be present, thus implying a non-linear optimization problem. This means that neither Eq. 1.9 nor Eq. 1.10 will have a closed-form solution; here, global optimization methods should be employed to find a good quality solution given the high amount of local optima (Jaqaman & Danuser, 2006; Koch, 2013; Saa & Nielsen, 2017). Several global optimizers have been developed for the resolution of this type of problems, which mainly differ in their algorithmic nature (i.e., stochastic or deterministic). The former, though not providing guarantee of global optimality, generally have a lower computational burden when used, thus representing an efficient alternative for solution space exploration. Deterministic methods, contrarily to the above, represent more robust optimal solutions, as they generally employ gradient-based solvers for local optimal exploration varying in how initial values are scattered; this, at the cost of a greater computational effort for exploration. Whichever method is selected, they cannot guarantee we are in presence of the global optimal, so that selection must be performed guided by the modeler's necessities.

c) A posteriori regression diagnostics

Once the maximum likelihood estimator ($\hat{\theta}$) is determined, *a posteriori* (posterior to regression) diagnostics must be performed to evaluate the model's validity. Post-regression diagnostics include the evaluation of several of robustness indices, where we can find parameter regression goodness-of-fit, parameter significance, and practical identifiability (or determinability) assessment.

i) Model goodness-of-fit

Goodness-of-fit is a relevant concept during robustness assessment, as it aims to determine if differences between simulated and experimental data are because of natural measurement noise and not from the model's inadequacy to fit the experimental data. To evaluate the adequacy of a model, usage of standard checks (i.e., distribution of residuals), as well as statistical testing can be implemented (Franceschini & Macchietto, 2008). The above generally consists in the use of null hypothesis testing, where the expected minimal value of Eq. 1.10 is tested to be equal to the number of degrees of freedom from model regression (Jaqaman & Danuser, 2006). Other criteria used for goodness-of-fit evaluation are the employment of likelihood ratio tests, as well as other similar criteria such as the Akaike or Bayesian Information Criterion, where smaller values imply a better overall goodness-of-fit in a model (Akaike, 1998; Jaqaman & Danuser, 2006; Lehmann & Romano, 2006; Saa & Nielsen, 2017).

ii) Parameter determinability

Parameter determinability is related to the presence of hidden interdependencies in the data, which precludes parameters to be uniquely determined. To evaluate this dimension of model robustness, interdependency between model parameters can be assessed through the variance-covariance matrix obtained during model regression. If values of this matrix are near extreme values (-1 or 1), then this reflects that two parameters are strongly influencing each other, thus compensating their values during regression. Overall, this situation tends to be detrimental to

a model's reliability, as this situation can lead to regression instability, which is the significant variation of estimated parameter values when adding new data for calibration; thus, the model would not be reliable, and the response would be inconsistent (Jaqaman & Danuser, 2006).

iii) Parameter significance

Finally, the statistical significance of each of the assessed model's parameters must be evaluated to corroborate if all parameters are significantly different from zero. The previous is to discard the possibility that a parameter turns out not affecting how independent and dependent variables relate in the model. A parameter can be characterized as non-significant because of several reasons: the large uncertainty within the data when compared to its output sensitivity in the model, lack of sufficient data, data not being informative, among others. To assess the parametric significance of estimates, the most popular tool is the Student *t-value*, an indicator reflecting how estimated values compare with their confidence intervals. Here, high values indicate parameters being reliable estimates, while lower values suggest the presence of the zero value in a parameter's confidence interval. The previous may also imply that some parameters are highly correlated, especially in models with many parameters and/or highly nonlinear structure (Franceschini & Macchietto, 2008).

1.3.2 Strategic reparametrization as a model enhancer

Usually, model calibration tends to be defective when hidden limitations in the data are not addressed by performing the previously stated tests. Normally, when models have an excessive number of parameters, these tend to adjust very well to the data used during the DPE process. However, these usually fail to represent the same modeled system when a new dataset is at hand, as well as when experimental conditions change even slightly. This is related to the high uncertainty in the

estimated parameters, which, when not assessed guided by the previous robustness-related concepts, lead to an unstable model in which predictions are biased (Egea et al., 2009; Jones et al., 2002). When the assessment of robustness is properly executed, there is a high chance of determining a defective parameter set in a given model structure, thus, adequate measures can be implemented to transform a problematic model into a more consistent variant. Following this idea, a commonly promoted strategy to correct problematic model structures is the reparametrization process, i.e., the generation of alternative nested models obtained from imposing linear constraints or fixing estimated parameters in a given model structure, so that the new model presents robustness indicators suggesting a superior model quality when compared to its initial structure (Krausch et al., 2019; Saa & Nielsen, 2017; Sánchez et al., 2014b). For example, if we suppose the hypothetical model in Figure 1-5 ($f_0(t, \theta_u) = t^{\theta_1 + \theta_2} + \theta_3$), which considers three estimated parameters ($\theta_u = [\theta_1, \theta_2, \theta_3]$), then a possible model structure originated by reparametrization could be $f_1(t, \theta_v) = t^{a + \theta_2} + \theta_3$, where θ_1 is replaced for a fixed value a , thus, the new parameter vector for this model structure is $\theta_v = [\theta_2, \theta_3]$.

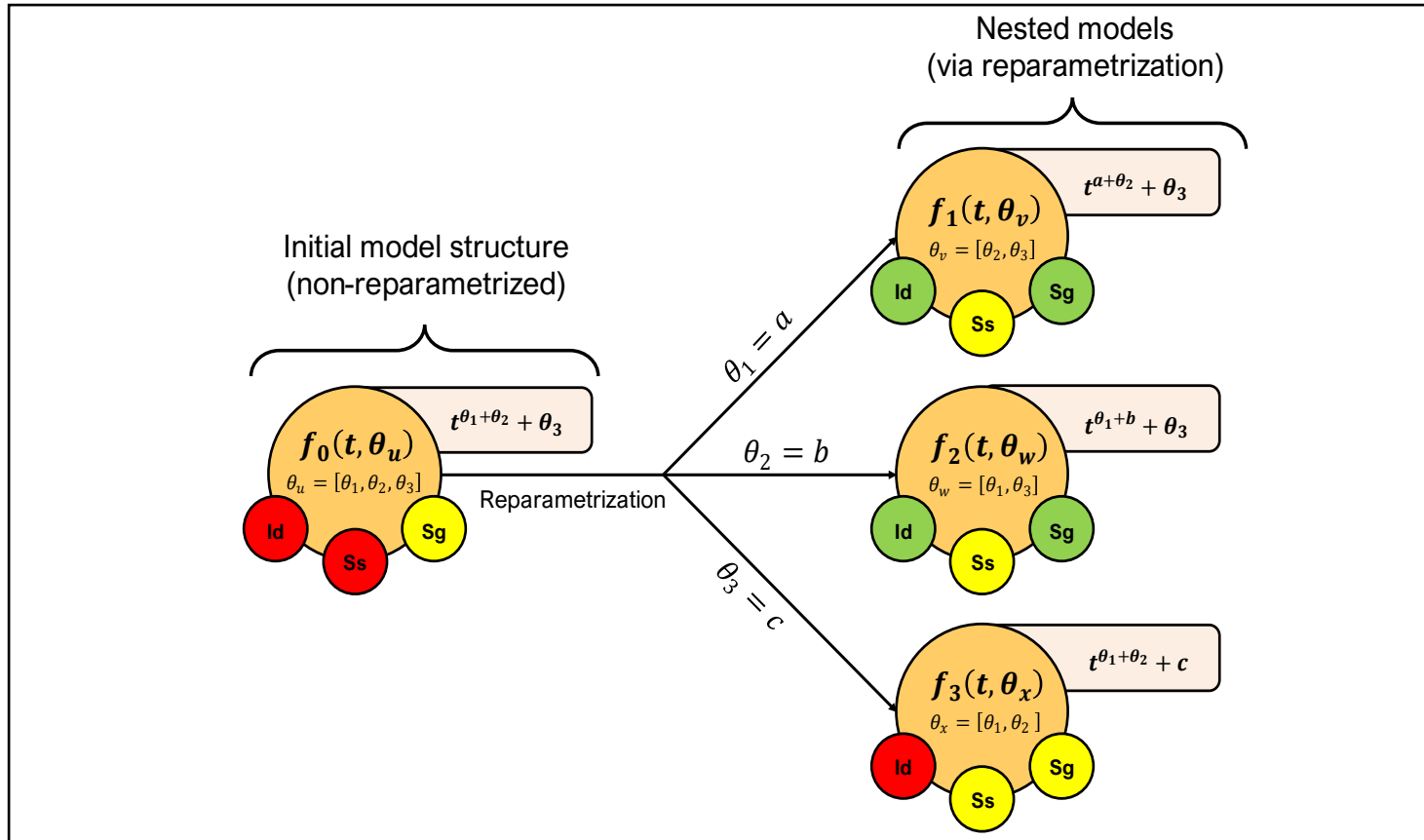


Figure 1-5. Graphic representation of the reparametrization process. Here, large orange circles represent a given model structure characterized by its unique parameters and associated identifiability (Id), sensitivity (Ss), and significance (Sg) performance indicators (small circles). Red, yellow, and green refer to problematic, acceptable, and robust features for each of the robustness dimensions.

In this example, both f_0 and f_1 are model structures, being the previous created from the original model structure (f_0) with its unique estimated parameter set (θ_v) and its own fixed values ($\theta_1 = a$). Here, the nested model structure f_1 presents different robustness indicators when compared with the original model structure f_0 , in this case being more robust f_1 , as parameters in the exponential $\theta_1 + \theta_2$ are clearly impossible to determine uniquely or are perfectly correlated.

The main purpose of the reparametrization process is to reduce the model complexity by fixing its problematic parameters so that the best model structure to fit the available data is achieved. Several authors have proposed diverse methods for the individual evaluation of the previously mentioned robustness dimensions (Jaquaman & Danuser, 2006; Kreutz et al., 2013; Saa & Nielsen, 2017; Stigter & Molenaar, 2015). An interesting method was proposed by Sánchez et. al (2014b), where they present a Heuristic Iterative Procedure for Parameter Optimization (HIPPO). The above corresponds to an algorithm for robustness assessment and DPE, where pre/post-regression diagnostics are deeply involved in its functioning. When given as input a model and experimental data, this algorithm generates as output reduced (or reparametrized) model structures with no problematic parameters, including the optimal values adjusting the estimated parameters of this model structure ($\hat{\theta}$). Given parameter interdependency and multiplicity of defective parameters, HIPPO uses an iterative method to evaluate model structures, where if a model structure presents one or more problematic parameters, one of these are fixed, and the process is repeated until no problematic features are detected.

1.4 Approach of the thesis

While several methods for robustness assessment exist, these tend to solely focus on the evaluation of a single dimension of the previously referred robustness dimensions. The HIPPO algorithm is the exception of the above, as it manages to

integrally assess robustness while performing DPE, yielding simpler models with no problematic features (Sánchez et al., 2014b). However, this last method is limited, as reparametrization is reduced to the selection of the first model structure with no problematic features. In this last approach, two problems can be addressed. First, robustness is treated in a binary way; this means, we seek to answer the question: ¿Is this model structure robust? However, a different approach could be answering the questions: ¿How robust is this model structure? ¿How sensitive is it? ¿How certain are we about this robust model's parameters? The questions above assume robustness as a continuous property of models, where each dimension (significance, identifiability, sensitivity, and goodness-of-fit) is measured quantitatively. This assumption is convenient, as a strategic approach for comparing non-problematic robust models can be performed guided by these concepts. Secondly, there are no clear methodologies when rating a model's robustness, i.e., it is not clear if, for example, a robust model structure with better sensitivity attributes is preferred over another one with higher certainty over its estimated parameters.

Considering the above, the purpose of this thesis is to approach the previously stated difficulties in a strategical way, validating and implementing results as part of the enhancement of FPM models used in the SmartWinery platform. Following this line, the following subsections intend to concisely inform readers about the main and secondary objectives, as well as the methodology employed to accomplish the above. Also, the main results and conclusions from this work will be briefly reviewed. Closing this section ends Chapter I, opening into Chapter II, which shows a submitted journal manuscript which summarizes the work performed to accomplish the objectives stated in the following subsections.

1.4.1 Hypothesis and objectives

a) Hypothesis

It is possible to systematically select a model structure which best enables to guarantee predictive quality by generating a workflow that strategically implements robustness assessment algorithms combined with multicriteria decision-making algorithms during the dynamic parameter estimation procedure with limited data.

b) Main objective

The main objective of this work is the generation of a reparametrization strategy that enables to obtain reparametrized AF models via robustness assessment methods, yielding the best performing model when data is limited by combining the use of robustness indicators with multicriteria decision-making methods.

c) Secondary objectives

To accomplish the previously stated objective, a series of secondary objectives were defined, which are the following:

- i) Generate data associated with the fermentation kinetics by conducting experimental vinifications at different process scales (laboratory and pilot scales).
- ii) Use of data generated in laboratory-scale fermentations for model processing in the HIPPO algorithm (Sánchez et al., 2014b) to characterize, evaluate, and select the model reparametrization that best represents the data through multicriteria decision-making methods.

- iii) Apply the selected best performing AF model structure obtained in the previous analysis to validate its reliability through simulations of the pilot-scale vinification experiments.

1.4.2 General methodology

A series of fermentations were carried out simultaneously in batch reactors of 5 (laboratory-scale) and 1000 (pilot-scale) liters at CRI; this, following a standardized industrial winemaking protocol to avoid processing discrepancies. During fermentation, samples were taken periodically and analyzed using an automatized spectrophotometer (Y15, Biosystems), where measurements of the main fermentation kinetics-related species (glucose, fructose, and yeast available nitrogen) were measured for data collection. Complementing the above, operational data consisting in must and cap temperature was measured through PT100 sensors installed in the fermenters. Once experimentation was fulfilled, laboratory-scale data was implemented into the HIPPO algorithm to generate reparametrized model structures with their corresponding robustness indicators (from pre/post regression diagnostics) when characterizing the Zenteno et al. (2010) and Coleman et al. (2007) wine fermentation models. Posteriorly, model structures were filtered by establishing robustness viability thresholds, and consecutively best-performing robust models were selected by applying a series of multicriteria decision-making methods (Wang & Rangaiah, 2017). Finally, validation of the best performing model structure selected from the previous analysis was assessed using the pilot-scale experimental data and simulations, where robustness indicators, as well as prediction performance, were analyzed and contrasted to the original model's performance.

1.4.3 Principal results and conclusions

a) Preliminary testing of the analyzed models

To establish a comparative point before applying the method proposed in this thesis, two wine fermentation models (Zenteno and Coleman) were calibrated and tested to evaluate their performance in prediction generation over an independent set of validation experiments in both laboratory and pilot scales. Here, it was seen that when applied for laboratory-scale experiments, both calibrated models showed an overall good predictive performance, observing an averaged global performance index (GPI) of 0.93 between all model states. This was not the case for pilot scale experiments, as samples from these were less frequent and contained lower quality information for parameter calibration, leading to lower values respecting the fitting performance (average GPI of 0.57).

b) Selection and testing of the best model structure selected by the procedure

Conclusions from the above supported the idea of reparametrizing the Zenteno and Coleman models to select a model structure which could best use the limited data available for calibration to enhance predictive performance. Here, the method developed in this work was applied in the previous models, and reparametrizations (i.e., model structures originated from the initial Zenteno and Coleman models, but differing in free and fixed parameters) were compared to select the best performing ones over the laboratory scale experiments. The best performing model structure was defined as “Zenteno-10325”, presenting the original structure of the Zenteno model while having 6 free and 8 fixed parameters. When applied to simulate laboratory scale validation experiments, Zenteno-10325 showed a superior performance compared to the original models, increasing averaged GPI from 0.93 to 0.97.

c) Model validation using pilot-scale data

Model structure Zenteno-10325 was further tested via its application in the pilot scale validation experiments. Again, an increase in predictive performance was observed when compared with original model performance, significantly increasing fitting performance in sugar-related model states by around 22% of the original model's averaged values.

d) Limitations and future perspectives

Though reflecting successful results, work must be done to establish this method as a viable one. Some activities include the further validation of this model identification method by supporting results with experiments varying in initial and operational conditions; this, to quantify their effect in the parameter estimation process, giving more information about the effectiveness of the proposed method in the mission of generating reliable model structures. Moreover, another concerning problem is the high computational cost incurred by some of the methods implied in this procedure. It is discussed that the HIPPO algorithm took around 13 days to be executed in both Zenteno and Coleman models. This situation is heavily detrimental if a wide application of this method is to be expected, as computational time grows exponentially as parameters increase in number; thus, this situation needs to be addressed urgently so that viability is insured. Disregarding the above, the work in this thesis proved that a reparametrization strategy can lead into parameter estimates and model structures that predict better when compared to models not presenting a reparametrization considering the available calibration data. It is expected in the future an overall enhancement of this method, positioning it as a generic valuable tool in modelist's toolboxes when working towards model calibration when limited data is at hand.

2. ENHANCING WINE FERMENTATION MODELS IN THE PRESENCE OF LIMITED DATA: A MULTI-CRITERIA GUIDED PARAMETRIC ROBUSTNESS ASSESSMENT WORKFLOW

2.1. Abstract

Climate change has shortened vintage periods, complicating the management of fermentation capacity. Technological paradigms, such as model-based design, industry 4.0, and the internet of things, can play a crucial role in simplifying the fermentation process's operation and management. Oenological decision-making can be significantly improved and simplified with fermentation models applicable in real-time. Bioprocess-related models like these typically include many unknown parameters where, given the high cost they incur, must be estimated from scarce and noisy data obtained from small-scale experiments. This leads to unreliable estimations and overfitting, especially when process scaling is at hand. This study developed and applied a robust model assessment workflow to reparametrize and select wine fermentation models to enhance the parameter estimation process. Several computational tools were integrated to identify a model structure that can be fitted with limited data, while presenting no parametric identifiability, sensitivity, and significance problems. Model selection employed several multi-criteria decision-making algorithms considering robustness criteria and statistical assessments. To validate the obtained results, wine fermentations were carried on 5L (laboratory-scale) and 1000L (pilot-scale) reactors, defining calibration and validation experimental sets for each case. Using the laboratory-scale data, and guided by robustness and performance indicators, the model structure which optimally adapted to the limited data was selected (named Zenteno-10325), increasing the non-reparametrized model's global performance index from 0.92 to 0.97. To further validate the application of this method for process scaling, the Zenteno-10325 model structure was applied to simulate pilot-scale experiments, leading to a 24\% increase over the sugar state prediction adjustment determination coefficient. Overall, this methodology generated better performing reparametrized structures when compared with non-reparametrized

ones. Though applied to fermentation models, this method can be extended to other systems and applications like model-based experimental design.

2.2. Introduction

As a consequence of climate change, the harvest season's behavior in vineyards has become increasingly unpredictable (Gurbey, 2020). This unpredictability has affected the adequate management of harvest dates, truck scheduling, grapes reception, and winery capacity (Webb et al., 2007). Winemakers dispose of a broad set of operational options and technologies to ensure that wine is made on schedule and with the expected quality (Mira de Orduña, 2010). Model-based methods are becoming increasingly popular in the food industry to achieve better and more consistent operations since they provide a faster and more reliable solution than experimental trial and error (Bordons & Núñez-Reyes, 2008; Qin & Badgwell, 2003). In wineries, dynamic fermentation models can help winemakers decide the best operating strategy to apply in a given situation (Ribéreau-Gayon et al., 2006), leading to more effective, efficient, and sustainable processes. This approach is particularly relevant for industrial winemaking, representing one of the most energy-consuming sectors in the food industry (Galitsky et al., 2005). Assessing dynamic fermentation models' reliability and predictive capacity is imperative to guarantee the winemaking process's expected improvements.

Several alcoholic fermentation models have been proposed that differ on the included phenomena (Miller & Block, 2020). Boulton, (1980) described a simple dynamic wine fermentation model that considered sugar as the limiting substrate and that temperature affected the yeast specific growth rate only. To represent industrial wine fermentations better, Cramer et al. (2002) and Coleman et al. (2007) considered nitrogen as the primary limiting substrate and that many more model parameters were temperature-dependent. As oxygen has demonstrated a critical role during initial yeast development, some authors have incorporated its

effect on fermentation models as well (Saa et al., 2012; Cerda-Drago et al., 2016). Sainz et al. (2003), Pizarro et al. (2007), and Vargas et al. (2011) developed hybrid models that coupled a steady-state stoichiometric metabolic model (representing the internal yeast fluxes) with biomass and metabolites dynamic mass balances (dynamic flux balance analysis). Most of the above models can provide reasonable predictions of the principal metabolites' dynamics at both laboratory and pilot-scale fermentations. However, these model predictions are less reliable at an industrial scale due to the inevitable temperature and concentrations gradients, especially in red wine fermentations (Miller & Block, 2020). At this scale, compartmental (Zenteno et al., 2010) or computational fluid dynamic models are needed (Miller et al., 2019).

Once an adequate model is selected, dynamic model parameter estimation (model calibration) must be performed, which is challenging for bioprocess models typically characterized by strong non-linearities, many fitting parameters, as well as scarce and noisy data (Bonate, 2011; Sacher et al., 2011; Krausch et al., 2019). It is advisable to carefully fix a priori some fitting parameters based on previous knowledge or preliminary estimations in these cases. Therefore, only a subset of the original model parameters are estimated to different data sets, yielding a reliable and robust predictive model; this process is called reparametrization. A model structure can be qualified as robust when predictions are reasonably accurate (predictive capacity) and when all estimated parameters affect model predictions significantly (sensitivity), are estimated with reasonable precision (identifiability), and are significantly different from zero (significance). Nevertheless, adequate reparametrization is hard to achieve, requiring several regression diagnostic techniques to assess the iterative parameter fixing procedure (Saa & Nielsen, 2017). These diagnostic techniques can be classified as a priori and posteriori. A priori techniques are frequently applied for structural identifiability assessments (Jaqaman & Danuser, 2006; Sacher et al., 2011; Stigter & Molenaar, 2015). A posteriori methods refer to those that evaluate model

calibration results, assessing, for example, predictive capacity, parsimony, parametric identifiability, and parametric significance (Jaqaman & Danuser, 2006; Saa & Nielsen, 2017). A wide diversity of statistical techniques have been proposed to assess either specific aspects (Jaqaman & Danuser, 2006; Rodríguez-Fernández et al., 2007; Kreutz et al., 2013; Stigter & Molenaar, 2015; Saa & Nielsen, 2017) or overall performance (Sánchez et al., 2014b). Despite its importance for model reliability, regression analysis is not widely applied in bioprocess modeling. Typically, dynamic process models are used without an exhaustive robustness analysis (Streif et al., 2013), leading to local instability and overfitting that compromise the models' predictive performance when used under different experimental conditions (Jaqaman & Danuser, 2006).

This study proposes a novel workflow to select a robust dynamic model structure for fitting wine fermentations. The method was applied to select suitable model structures for two dynamic models (Coleman et al., 2007; Zenteno et al., 2010), using laboratory and pilot-scale wine fermentation data. The proposed workflow initiates with the iterative reparameterization of the evaluated models using the HIPPO package (Sánchez et al., 2014b), where several model structures were generated by fixing different parameter sets. Then, several regression performance indices were calculated for each model structure generated. According to the regression and global performance indices, several multi-criteria decision-making (MCDM) algorithms were applied to choose the most suitable model structure, validating the selection with laboratory and pilot-scale data.

2.3. Materials and Methods

2.3.1. Experimental design

A total of 6 experimental fermentations were performed using *Cabernet Sauvignon* grapes obtained from different producers scattered throughout the Maule region,

Chile. These fermentations were carried out simultaneously at two scales, referred to as *laboratory-scale* and *pilot-scale* experiments. Each fermentation system, the fermentation procedure, and the sampling methods are described in the following sections:

a) *Laboratory-scale fermentations*

Laboratory-scale experiments were carried out on a 5L reactor, where the temperature was maintained at 26 °C by a heating/cooling jacket covering around 70\% of the reactor's wall (Fig. 2-1a). As typically required in red wine production to enhance pomace extraction and nutrient distribution, pumping-over was carried out three times per day for two minutes at a 2.5 L min⁻¹ flow rate by a pump-powered aspersion system.

b) *Pilot-scale fermentations*

Pilot-scale experiments were carried out on a 1000L cubic reactor, where the must temperature was controlled at 26 °C by a heating/cooling coil placed at the center of the reactor (Fig. 2-1b). Pumping-over was carried out using the same aspersion system and operating parameters described above, scaled-up to the corresponding size.

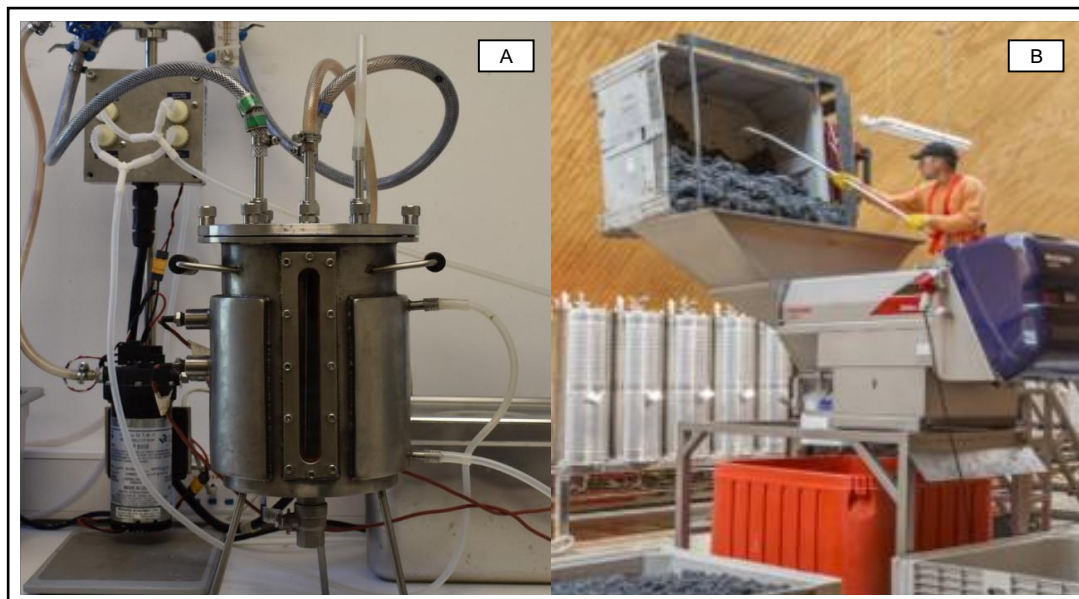


Figure 2-1. Laboratory and pilot scale batch experimental setups used for the experiments. **A.** Laboratory scale 5 L reactor. **B.** Pilot scale 1000 L reactor.

c) *Red wine fermentation and juice preparation*

For all fermentation runs, grapes were crushed, obtaining 70% of juice and 30% of grape solids (skins and seeds). The pre-fermentative juice was corrected to 23.5 °Brix, pH between 3.4-3.5, and yeast assimilable nitrogen (YAN) around 250 mg L⁻¹. This correction was achieved by diluting with distilled water, and by adding tartaric acid, diammonium phosphate (DAP), and sulfur dioxide (5 g hL⁻¹) through a potassium metabisulfite solution (50 g L⁻¹). Each fermentation was inoculated with the same *Saccharomyces cerevisiae* yeast strain (Maurivin PDM, ENSIS Sciences) at a concentration of 20g hL⁻¹. Also, DAP was added in the middle of the fermentation (density 1050 g L⁻¹) to avoid sluggish or stuck fermentations.

d) *Sampling and analysis procedures*

A total of 13 and 9 samples were taken to measure glucose, fructose, and YAN concentrations in each *laboratory-scale* and *pilot-scale* experiment, respectively. An automatic spectrophotometric analyzer (Y15, BioSystems) with specific

reactive kits was used for sample analysis. Sugar consumption (°Brix and density) was measured with a portable densimeter (DMA35, Anton Paar). Must and cap temperatures were measured with PT100 sensors and recorded with a sampling frequency of 1 minute.

2.3.2. Workflow methodology

Fig. 2-2 shows a general scheme of the robust parameter estimation workflow applied in this work. The method begins with the application of the Heuristic Iterative Procedure for Parameter Optimization (HIPPO) algorithm (Sánchez et al., 2014b) to characterize many model structures (typically in the order of thousands) with different sets of free and fixed model parameters (reparametrization). Here, a single calibration experiment is used to analyze the model structures. A priori and a posteriori regression diagnostic techniques are applied to all the generated models in this analysis. A priori methods focus solely on evaluating structural identifiability, while a posteriori techniques assess several dimensions of model robustness: goodness-of-fit, parameter determinability (or practical identifiability), and parameter significance. Different criteria were applied in the second stage to choose a reduced set of models (typically less than 50) with desirable characteristics. These selected model structures are further assessed using several multi-criteria decision-making methods, and a few of them (around 3) are pre-selected for the next stage. Finally, other statistical indices were applied to these models, with additional calibration and validation data (in our case, the rest of the laboratory data), ending with a single model. The final model is finally assessed with independent data (in our case, calibration and validation pilot-scale data). The workflow's final product is a minimum model structure that can fit data from different experiments of the same or similar system (in our case, laboratory-scale and pilot-scale winemaking fermentations). Details of each step of this workflow are provided in the following sections.

a) *Characterization of model structures using HIPPO*

The first step in this workflow is to generate model reparametrizations using the HIPPO algorithm. Each model structure was defined by the same model equations but with different combinations of free and fixed parameters. Model structures are calibrated by least-squares regression using the meta-heuristic optimization code Enhanced Scatter Search (eSS, Egea et al., 2014) in HIPPO. After calibration, a priori and a posteriori regression diagnostics were applied to assess parameter identifiability, significance, sensitivity, and fitting-performance of each model structure. Parameter identifiability aims at deciding if a model's parameters can be estimated uniquely considering structural identifiability and practical identifiability. Structural identifiability, computed before parameter fitting, depends on how sensitive are the model outputs to the free (fitting) parameters of a given model structure. HIPPO computes a normalized sensitivity score matrix (SSM) for each model structure using the method described in Hao et al. (2006). Practical identifiability, computed after parameter fitting, measures the capacity to estimate the free parameters' true values using a given model structure and a given data set. Using the regression results, HIPPO generates a parameter correlation matrix (PCM) applying the standard Pearson correlation formula to determine highly correlated parameter pairs.

Parametric significance determines if the free parameters of a given model structure are significantly different from zero. Identifiable parameters can be insignificant due to the high uncertainty in the available data (Jaqaman & Danuser, 2006). HIPPO applies the *t-value* to detect insignificant free parameters,

$$t_{value}^p = \frac{4\hat{\theta}_p}{CI_p} \quad (2.1)$$

where $\hat{\theta}_p$ is the estimated value of parameter p , and CI_p is its confidence interval.

Typically, more than one model can fit the data well, and those with more free parameters tend to fit the data better, but at the cost of increasing the parameter interdependency (thus leading to identifiability problems). Hence, simpler and more parsimonious models that fit the data well are preferred (Jaqaman & Danuser, 2006). HIPPO calculates the Corrected Akaike Information Criterion (AIC_c), which is useful to measure the trade-off between model accuracy and the number of free parameters. HIPPO requires a specific data set and an ordinary differential equations (ODE) model with many fitting parameters. In our example, calibration data comprises temperature, fructose, glucose, and YAN measurements of one lab-scale fermentation (identified as LAB-LO(02)). Two wine fermentation models (Coleman et al., 2007; Zenteno et al., 2010) were analyzed, one comprising 4 ODEs and 12 fitting parameters, and the other comprising 5 ODEs and 14 fitting parameters (CO₂, density, and temperature ODEs not considered, refer to Appendix A).

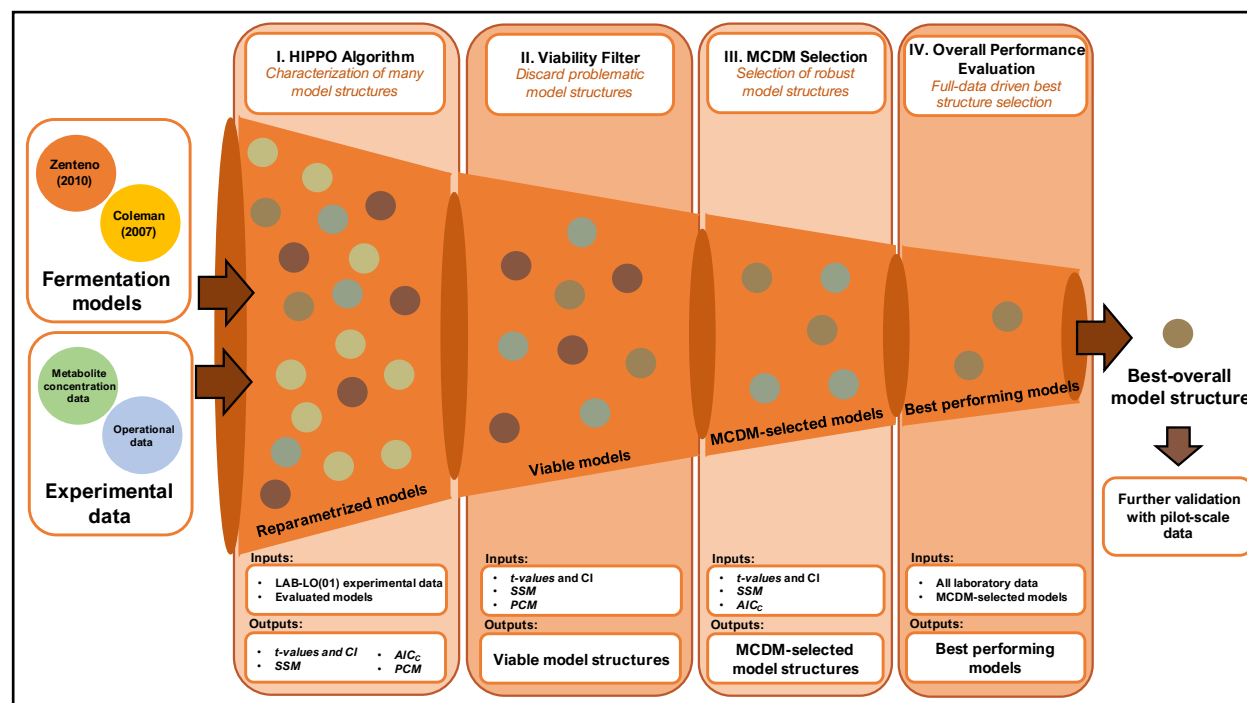


Figure 2-2. Graphical representation of the workflow used for model robustness assessment and selection. In step I the characterization of model structures is performed through reparametrization of the model used as input with a reduced set of the experimental data. Steps II and III refer to selection steps, where respectively viability thresholds and MCDM selections that were defined previously are used to evaluate the model structure's robustness indices. Finally, in step IV, model structures are recalibrated using all calibration data, and are then evaluated using a predefined validation dataset. The selected model in this final step is called the best-overall model structure, which is characterized by its highly robust features.

b) *Selection of viable model structures*

The viability analysis is the second step in the workflow, selecting a reduced set (typically less than 50) of the many reparametrized model structures generated by HIPPO. Here, the SSM (Hao et al., 2006) and normalized sensitivity curves (NSC) are re-calculated using the symbolic procedure described by Stigter & Molenaar (2015). The re-calculated SSM and the analysis provided by HIPPO (*t-values*, CI's, PCM, AIC_c) were considered for selecting the viable models according to the following criteria:

i) Parametric identifiability condition.

The cross-correlation among all parameter pairs of a given model structure is calculated (Jaqaman & Danuser, 2006). A threshold of 0.95 was established for the absolute values of non-diagonal elements of the PCM. Higher values indicate a high correlation between pairs of parameters, i.e., the structure is not identifiable.

ii) Significance condition.

t-values were evaluated as proposed by Sánchez et al. (2014b) to determine the statistical significance of parameters. The previous values were calculated for each of the p estimated parameters of a structure by using the nominal parameter values ($\hat{\theta}_p$) and their estimated confidence intervals (CI_p) as in Eq. 2.1. Values of this index lower than 2 suggest that zero is within the corresponding parameter's confidence interval. An adequate model structure should have all *t-values* > 2 .

iii) Sensitivity condition.

Here, the model structure's sensitivity was assessed using the re-calculated SSM. Each element of this matrix $SSM_{i,p}$ represented the

relative sensitivity score of the state variable i for the estimated parameter p . Following Sánchez et al. (2014b), we established a threshold value of 0.01 for all elements of the relative sensitivity matrix to qualify model structures as sensitive to their parameters.

c) *Selection of robust model structures*

Once the viable structures are chosen, the next step in the workflow consists of applying MCDM methods to obtain a few model structures (typically around 3). Here, the previously described estimation indices define O different criteria (OF_o), each reflecting a model's performance over a given robustness measure. To broaden the scope of this analysis, different weights (ω_o) for each criteria define different scenarios, and several MCDM methods were applied as proposed in Wang & Rangaiah (2017) to choose the best model structure for a given scenario, as shown in Table 2-1.

The evaluated criteria were designed to assess the different dimensions that characterize the reliability of a given model structure, i.e., identifiability, significance, sensitivity and goodness-of-fit. This is important since the performance in each case vary significantly; normally, a good performance in one of these dimensions is associated with a poor performance in other dimensions (e.g., a given model structure fits the data accurately, but the model is not sensitive to some free parameters). The model selection procedure makes use of the output indices from HIPPO ($\hat{\theta}$, CI and AIC_c) and the symbolically calculated NSC as follows:

i) Performance criteria.

The corrected Akaike Information Criterion (AIC_c), provided by HIPPO and calculated like in Eq. 2.2, assess the fitting performance. It considers the square sum of errors (SSE), the number of estimated parameters (P),

and the number of experimental points (N). The lower the AIC_c value, the better the model performance is (should be minimized).

$$AIC_c = N \cdot \log\left(\frac{SSE}{N}\right) + 2(P + 1) + 2 \frac{(P+1)(P+2)}{N-P-2} \quad (2.2)$$

ii) Significance criteria.

This criterion, designed to measure the overall significance of a model structure, is given by the mean of the normalized confidence intervals (MNCI) of the P parameters provided by HIPPO (Eq. 2.3). Small values indicate a better overall significance of the model structure (should be minimized).

$$MNCI = \frac{1}{P} \sum_{p=1}^P \frac{CI_p}{\hat{\theta}_p} \quad (2.3)$$

iii) Sensitivity criteria.

Total parametric sensitivity ($G_{i,p}$) was calculated as the integral of the symbolically calculated NSC for each i state variable for each p estimated parameter, as follows:

$$G_{i,p} = \int_{t_0}^{t_f} NSC_{i,p} dt = \int_{t_0}^{t_f} \frac{1}{\max(\hat{x}_i)} \frac{dx_i}{d\theta_p} dt \quad (2.4)$$

where $\max(\hat{x}_i)$ is the maximum predicted value of the i model state, $\frac{dx_i}{d\theta_p}$ is the parametric sensitivity of each state with respect to each parameter, and t_0 and t_f refer to the initial and final integration times, respectively. Then, the global parametric sensitivity score (GS_i) for each state i was calculated as follows:

$$GS_i = \sum_{p=1}^P G_{i,p} \quad (2.5)$$

where each GS_i is a sensitivity-related measurement associated to each model state. In the MCDM assessment, each sensitivity measurement (GS_i) was weighted the same. Higher values of these indices indicate higher model sensitivity (should be maximized).

Table 2-1. MCDM methods and evaluated scenarios involved in the viable structure selection stage.

MCDM Method	¿Weights?	Principle	
TOPSIS	Yes	Euclidean distance to ideal	
LINMAP	Yes	Euclidean distance to ideal	
VIKOR	Yes	Euclidean distance to ideal	
SAW	Yes	Simple weighted summation	
MEW	Yes	Multiplicative weighting	
GRA	No	Similarity to utopic optimal	
FUCA	Yes	Global ranking	
Scenario Priority	Performance Criteria Weight	Significance Criteria Weight	Sensitivity Criteria Weight
No priority	0.33	0.33	0.33
Heuristic	0.50	0.30	0.20
Performance	0.80	0.10	0.10
Significance	0.10	0.80	0.10
Sensitivity	0.10	0.10	0.80

d) *Best-overall structure selection*

The last step of the workflow yields the best-overall model structure. Here, the few model structures obtained from the previous step are re-analyzed using the complete calibration and validation data sets (in our case, all laboratory-scale experiments in Table 2-2). Goodness-of-fit, and additional statistical testing and regression performance indicators, commonly used for model validation were applied (Bonate, 2011; Serebrinsky et al., 2019). Details of these procedures are given below.

Table 2-2. Experimental data sets.

Usage	Nº	Laboratory-scale	Pilot-scale
Calibration	1	LAB-LO(02)	CII-LO(07)
	2	LAB-LO(02)	CII-LO(08)
	3	LAB-MA(04)	CII-MA(09)
	4	LAB-MA(05)	CII-SR(15)
Validation	5	LAB-LO(07)	CII-MA(01)
	6	LAB-LO(08)	CII-LO(27)

i) MCDM models calibration

After MCDM selection was made, parameter estimation was performed over the selected structures using the laboratory-scale set of calibration experiments (Table 2-2) while applying a bootstrapping-based regression procedure. The aforementioned was used to obtain the overall model structure as a result of fitting laboratory scale measurements. In this

procedure, each element corresponds to a specific experimental data set (e.g LAB-LO(02)), and different subsets contain different combinations of elements. Given K calibration experiments, a total of $\binom{K}{v}$ subsets are formed for a determined subset size v . To simplify indexation, we identify subset $\Theta_{u,v}$ as a specific u subset with v elements. For example, [LAB-LO(02), LAB-LO(03), LAB-MA(04)] may correspond to subset $\Theta_{1,3}$, while [LAB-LO(02), LAB-LO(03)] and [LAB-LO(02), LAB-MA(04)] may correspond to subsets $\Theta_{1,2}$ and $\Theta_{2,2}$, respectively. Additionally, a squared-error function was defined considering the simulated ($x_{i,n}(\hat{\theta})$) and measured values ($x_{i,n}$) associated to each k -th experiment in a given subset. For each experiment, N_k measurement points were considered in the calculation of the squared-error over a total of I model states, as shown in Eq. 2.6.

$$J_k(\hat{\theta}) = \sum_n^{N_k} \sum_i^I \left(\frac{x_{i,n}(\hat{\theta}) - x_{i,n}}{\max(x_{i,n})} \right) \quad (2.6)$$

Finally, by combining Eq. 2.6 with the $\Theta_{u,v}$ subsets, the global error function $GE(\hat{\theta})$ for parameter calibration of the MCDM-selected models was defined as the sum of the errors calculated for all subsets (Eq. 2.7).

$$GE(\hat{\theta}) = \sum_{v=1}^V \sum_{u=1}^{U_v} \sum_{k \in \Theta_{u,v}} J_k(\hat{\theta}) \quad (2.7)$$

where K corresponds to the total amount of calibration experiments, and U_v corresponds to the number of possible combinations with v elements (equivalent to $\binom{K}{v}$).

The parameter estimation problem was solved using eSS with *fminsearch* for the local search. We used the Matlab Parallel Computing Toolbox to accelerate this process, given the high computational cost incurred.

ii) Global performance index

Predictions generated by each of the calibrated MCDM-selected structures were evaluated using the adjusted determination coefficients (R_{adj}) from each of the measured model states (in our case glucose, fructose, and YAN). We defined an index to consider the multiplicity of model states and validation experiments. Considering the existence of a total of Q validation experiments and I measured model states, we defined the global performance index of a model structure w over a validation experiment q ($GPI_{w,q}$), as in Eq. 2.8.

$$GPI_{w,q} = \frac{1}{I} \cdot \sum_{i=1}^I (R_{adj})_{i,q}^w, \quad q = 1 \dots Q \quad (2.8)$$

$GPI_{w,q}$ represents the average between the adjusted coefficients of determination $(R_{adj})_{i,q}^w$ calculated from simulations of the i measured model states obtained from a given calibrated MCDM-selected model and validation experiment.

iii) Residual autocorrelation and normality

Statistical testing was applied over the residuals obtained from each selected structure to determine if the resulting models were biased. Anderson-Darling test (AD) was applied for the residual normality distribution and the Durbin-Watson test (DW) for residuals correlation. Both tests were applied with a significance level of 5%. Critical values for the DW test were obtained from Farebrother (1980).

iv) Structural identifiability stability analysis

To further ensure results from the structural identifiability analysis performed by HIPPO and check visually regression stability, we employed the method proposed by Stigter & Molenaar (2015). Here, by computing the Relative Output Sensitivity Matrix (ROSM) using symbolic procedures for parametric sensitivity calculation in combination

with Monte Carlo simulation techniques, the authors derive a method to graphically evaluate parametric identifiability and instability. This is performed by repeatedly calculating the ROSM while varying parameter's nominal values in a predefined interval. In each iteration, values of the singular values and associated right singular vectors are obtained through Singular Value Decomposition (SVD) of the ROSM. Consecutively, these are evaluated to analyze identifiability problems, where higher values suggest a high correlation between a given set of parameters. We refer readers to Stigter and Molenaar (2015) for a further understanding of this method.

v) Best structure selection

The final step is the selection of the best-overall model structure using the performance indicators from the prior sections. Here, validation data (in our case, laboratory-scale validation experiments in Table 2-2) is used to compute $GPI_{w,q}$, AD, and DW indices for each calibrated MCDM-selected model structure. The best-overall model structure should present a significant sensitivity to calibration, a stable response to small changes in parameter values, a good predictive capacity, and minimum bias, according to the already defined performance indices and analyses. Therefore, the best-overall model structure presents the highest averaged $GPI_{w,q}$, and the highest number of normally distributed uncorrelated residuals.

2.3.3. Model validation

Commonly, models are used in a setting that differs from the one used to calibrate them. For example, in our case, we calibrate the model using lab-scale experiments, but we would like to use the model in a pilot-scale fermenter. Therefore, the best-overall model structure was calibrated with the calibration

subset of the pilot-scale experiments (Table 2-2). The pilot-scale validation data was then used to assess the model performance and compare it with the full model (all model parameters were considered free). Here, the indices described above (GPI, AD p-value, and DW d-value) were applied to compare them.

2.4. Results and discussion

2.4.1. Initial evaluation of Coleman and Zenteno models

To accomplish the validation of the proposed methodology, we initially assessed the predictive capacity of the original Coleman (RCM, Coleman et al., 2007) and Zenteno (RZM, Zenteno et al., 2010) model structures, i.e., considering the fixed and free parameters proposed by the corresponding authors. These model structures were re-calibrated at both studied scales to adapt them to our data following the methodology in section 2.3.2d. In our experiments, we observed a case-specific YAN concentration (YAN stagnation point) at which yeasts seemed to stop consuming nitrogen (between 15 and 60 mg L⁻¹); other authors have also observed this (Childs et al., 2015; Coleman et al., 2007). Also, Coleman and Zenteno's original models consider the total consumption of the available nitrogen and no additions during the fermentation; hence, our simulations and experimental data were rearranged to consider the YAN stagnation points and the DAP additions. At each experiment, the YAN stagnation point was defined as the YAN concentration measured just before the DAP addition near must density 1050 gL⁻¹ (point of the fermentation where we typically observed YAN stagnation, refer to section 2.3.1c for information about our fermentation protocol). Then, we subtracted the corresponding YAN stagnation concentration from the YAN measurements, simulating a total YAN consumption; negative concentrations resulting from the subtraction were approximated to zero. DAP additions at the middle of the fermentations were considered by splitting the fermentation simulation in two: before DAP addition and after DAP addition. The initial YAN

values of the second simulation stage were modified according to the amount of DAP added. Once simulations were completed, YAN values were corrected by adding the YAN stagnation concentration to the simulated nitrogen consumption curve. YAN measurements close to zero (commonly those at the end of the fermentations) were excluded from calibration and performance assessment since these measurements introduced bias. In Fig. 2-3 the YAN corrections described above were graphically represented.

The performance of both models is summarized in Table 2-3, calculated with their corresponding validation datasets (Table 2-2). Sugar and density showed good precision ($R_{adj} > 0.79$); however, YAN predictions at the pilot-scale were inaccurate, probably due to the few measurements taken at the beginning of the pilot-scale fermentations (Jaqaman & Danuser, 2006). The protocol in these experiments considered a constant sampling based on the density evolution; samples were taken when must's density decreased by 10 kg/m^3 . This method resulted in 3 or less YAN measurements during the first 30 hours of fermentation, where nitrogen is consumed much faster than sugar. The laboratory-scale experimental protocol considered two measurements per day, independent of the density evolution, resulting in better model performance.

Table 2-3. Predictive performance over *laboratory-scale* and *pilot-scale* experimental data when simulating using the original Zenteno and Coleman models.

Coleman Model				
Process Scale	R_{adj} YAN	R_{adj} Sugar	R_{adj} Density	GPI
Laboratory	0.92 ± 0.02	0.93 ± 0.05	0.94 ± 0.02	0.93 ± 0.01
Pilot	0.22 ± 0.96	0.81 ± 0.15	0.81 ± 0.14	0.62 ± 0.35
Zenteno Model				
Process Scale	R_{adj} YAN	R_{adj} Sugar	R_{adj} Density	GPI
Laboratory	0.96 ± 0.04	0.97 ± 0.01	0.97 ± 0.01	0.92 ± 0.07
Pilot	-0.51 ± 1.61	0.79 ± 0.20	0.79 ± 0.19	0.52 ± 0.64

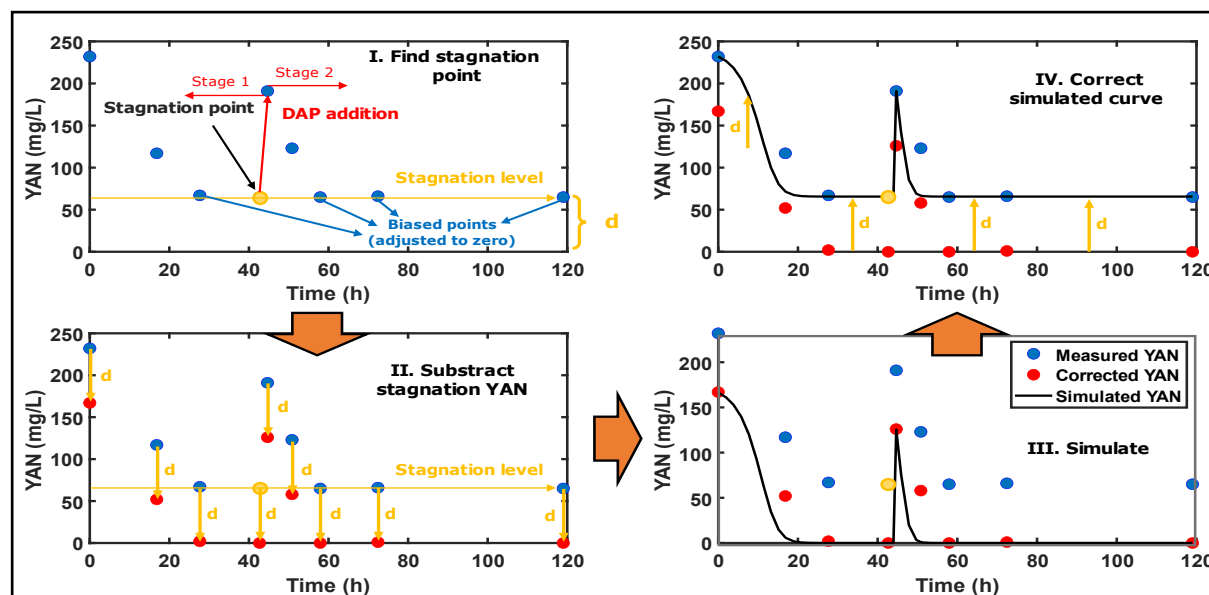


Figure 2-3. The four steps of the YAN correction procedure. Blue and red dots represent measured and corrected YAN measurements, respectively; continuous black line represents simulations. Step I identifies the stagnation point d as the YAN measurement before DAP addition. Measurements equal or lower than d are discarded and set to zero, and the fermentation is simulated in two stages (before and after DAP addition). In step II, YAN measurements are subtracted by d . In step III, simulations are carried out assuming corrected YAN values. Finally, in step IV, simulated values are corrected by adding the stagnation YAN concentration.

2.4.2. Best-overall model structure selection

In the first step of the proposed workflow (section 2.3.2), the reparametrized model structures Coleman et al. (2007) and Zenteno et al. (2010) are generated using the HIPPO algorithm (section 2.3.2a) and data from the calibration experiment LAB-LO(02). This computationally heavy procedure (around 13 days on an Intel 4770k processor) yielded 4092 and 11958 characterized Coleman and Zenteno model structures, respectively. The robustness and goodness-of-fit (section 2.3.2a) indices of the generated models were considered in steps II and III to select a reduced set of structures (Fig. 2-2). This selection was based on a single data set (LAB-LO(02)). A more reliable selection can be achieved using a more complete data set, although the computational burden would be much higher.

In step II, model structures with many limitations regarding identifiability, significance, and sensitivity were discarded according to the criteria specified in section 2.3.2b. The measurement noise and the discontinuous DAP additions in the middle of the fermentations hamper the ODEs integration and the parametric sensitivities calculation, which HIPPO cannot handle easily. Therefore, in this step, the SSM and NSC were recalculated for sensitivity assessment of all model structures generated by HIPPO using a more robust symbolic-based method (section 2.3.2b), although this method is computationally more demanding. This step yielded 29 and 5 “viable structures” of the Coleman and Zenteno models, respectively, which were processed further.

In step III, MCDM methods were applied (section 2.3.2c) to choose the most “voted” model structures under the five scenarios considered (Table 2-1). The decision-making results, considering the 5 viable Zenteno models, are shown in Fig. 2-4. This procedure allowed us to reduce the 29 Coleman viable models to just 3 model candidates (Fig. C-1) and yielded 2 Zenteno model candidates. The

candidate models were separated into two: the Coleman-MCDM candidate models (group I) and the Zenteno-MCDM candidate models (group II).

Table 2-4. Selected viable structure groups specifications.

Group	Fermentation model	Group GPI	Best GPI ^a	%Normal residuals	Uncorrelated residuals ^b (%)
Group I	Coleman	0.92 ± 0.03	0.94 ± 0.03	75%	50% (25%)
Group II	Zenteno	0.92 ± 0.07	0.97 ± 0.02	75%	62.5% (20%)

^a Denoted as the mean GPI obtained from simulations over the laboratory validation data set using the best performing model structure in each group.

^b Numbers appearing in parenthesis correspond to the percentage of residuals where correlation was uncertain in DW testing. These cases were not included in the calculation of the residual correlation percentage.

In step IV (Fig. 2-2), candidate models from groups I and II were re-calibrated with the four experimental calibration datasets (2.3.2.d). These re-calibrated model structures were validated with two additional laboratory-scale datasets (Table 2-2); these results are shown in Table 2-4. Both groups yielded the same average GPI (0.92), although group I presented a slightly lower standard deviation than group II, suggesting a more consistent predictive capacity. The best GPI was achieved by model 10325 in group II, which presented a better overall predictive capacity, considering both validation datasets. The best-performing (highest GPI) Zenteno-MCDM and Coleman-MCDM (3666) model structures were those that received most votes in at least one scenario by the MCDM algorithms (Fig. 2-4).

Sugar-related residuals (glucose, fructose, density, and total sugar) in both model groups typically followed a normal distribution (data not shown). In turn, nitrogen

was quickly consumed at the beginning of the experimental fermentations, and not many YAN measurements were taken, resulting in a few highly variable YAN residuals that did not follow a normal distribution.

The residuals of group II model structures were less autocorrelated than those in group I (Table 2-4). Hence, Zenteno structures provide a better description of the process dynamics, while Coleman structures cannot explain a significant part of the dynamic response, yielding less accurate predictions.

The original structures of both models (Coleman and Zenteno) were further tested with an alternative hybrid identifiability analysis method, which was proven as an efficient and reliable analysis to assesses the propagation of non-linear structural identifiability (Stigter & Molenaar, 2015). Performing 500 simulations assigning random values to each model parameter within the range $[\hat{\theta} \pm 0.2\hat{\theta}]$, we were able to detect the unidentifiable parameters of the original models; as nominal parameter values ($\hat{\theta}$), we used those generated by the HIPPO algorithm in the calibration of the original model structures.

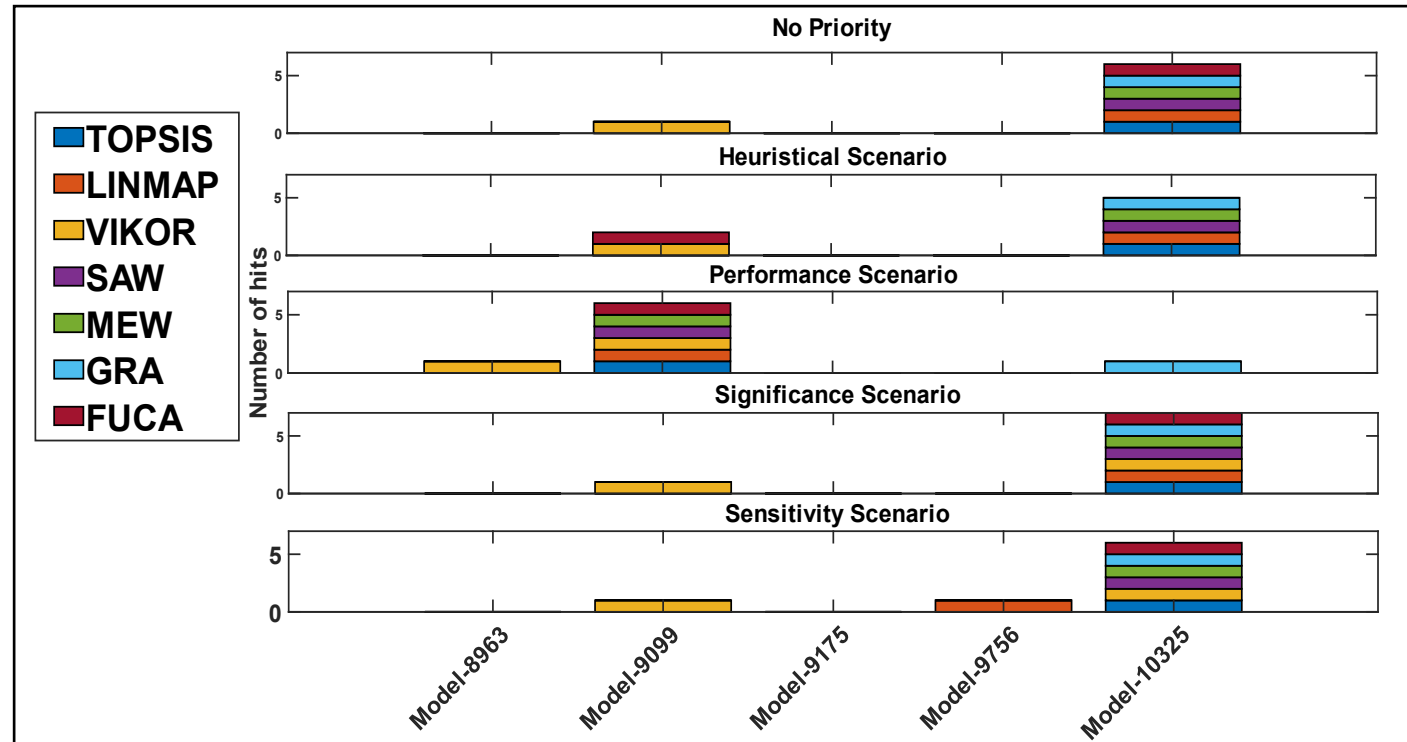


Figure 2-4. MCDM results when applied to model structures from group II (Zenteno originated models). Stacked bars represent votes committed by each of the MCDM methods listed in the figure legend (Table 2-1). Overall, model structures 9099 and 10325 represent good levels of robustness, as both models were the most voted in one or more scenario.

The Zenteno model results are shown in Fig. 2-5, while those of the Coleman model are shown in Figure D-1. In these figures we show the absolute values of components associated with the last singular vector from the SVD of the ROSM; averaged given the multiplicity of simulations. Those components that are significantly different from zero pinpoint the non-identifiable parameters, which for the Zenteno model are θ_{11} and θ_{12} (Y_{XG} and Y_{XF} , respectively; refer to Appendix A and Appendix B for further details about these parameters). These results showed concordance with the work leading into structure Zenteno-10325, as parameters tagged as problematic coincided with those fixed using our method. However, as seen in Appendix D, this analysis gave inconclusive results for the Coleman model, which were observed as instabilities when slightly changing parameter values, resulting in high standard deviations when applying this procedure. This indicates strong correlations between several of the parameters in this model; thus, suggesting inadequacy in the use of this model for our purposes.

Although this alternative analysis is limited to detect structural unidentifiability, it is much more robust than the HIPPO algorithms. Also, as HIPPO's algorithms intend to explore robustness following a tree-exploration approach (fixing a parameter implies the generation of a new exploration branch), those alternative methods could drastically benefit HIPPO's efficient implementation; these aim to allow early detection of problematic parameters. Consequently, it is advisable to run this hybrid algorithm before applying HIPPO, reducing the number of model structures generated significantly. This parameter reduction is especially relevant for metabolic genome-scale models, where the number of equations can be in the order of hundreds (Saitua et al., 2017; Sánchez et al., 2014a).

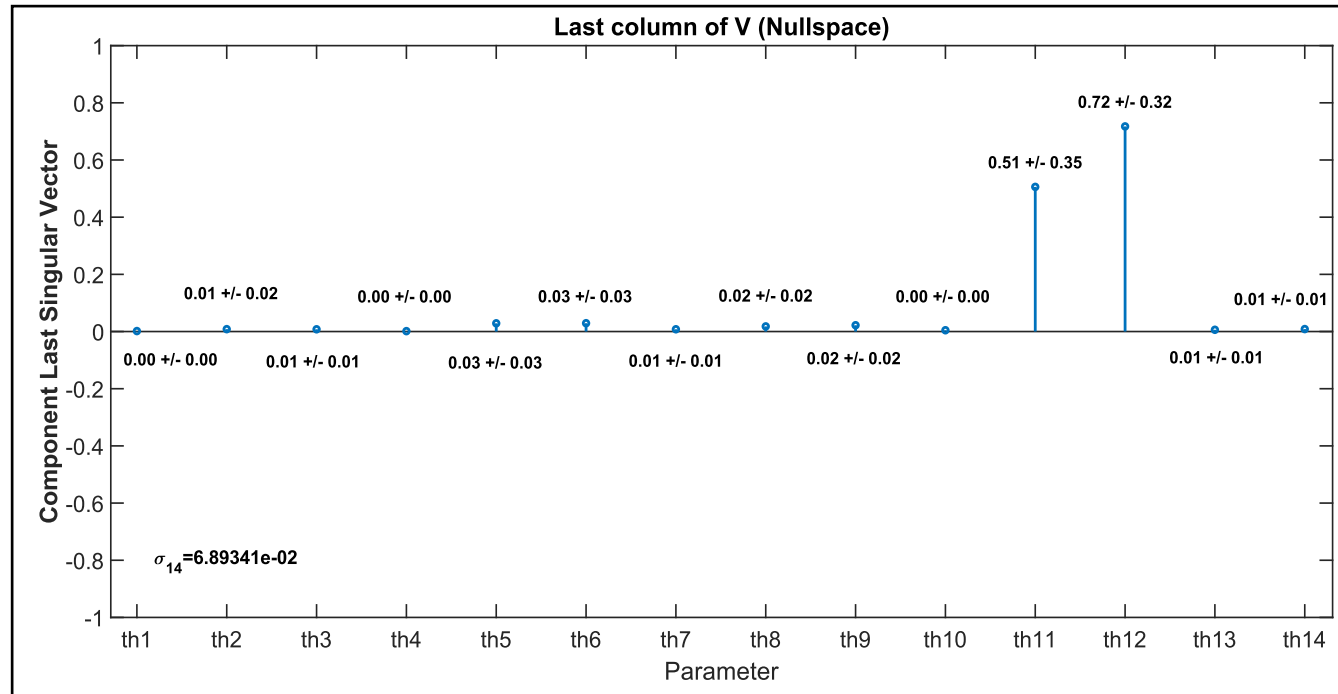


Figure 2-5. Mean absolute values of the components of the last singular vector of the ROSM obtained from applying the method proposed by Stigter & Molenaar (2015) over the reduced Zenteno model (Appendix A). A total of 500 iterations were performed, varying nominal parameter values. Values distant from zero represent problematic parameters given the data structure used for calibration (in this case, parameters θ_{11} and θ_{12} , refer to Appendix B). σ_{14} corresponds to the singular value associated to the last column in the SVD decomposition of the ROSM.

After applying the methods in step IV, the Zenteno-10325 model structure was selected for further assessment, given its superior performance. Parameter values of fixed and free parameters associated with this structure can be found in Appendix B.

2.4.3. Pilot-scale validation

Zenteno-10325 model structure was re-calibrated using the pilot-scale calibration data set (Table 2-2) to represent the fermentation's behavior adequately on this scale (see section 2.3.1b). We used the same calibration method described in section 2.3.2d. The values of the fixed parameters were those calculated by HIPPO (section 2.3.2a).

Predictive simulations were performed to reproduce the pilot-scale validation experiments (Table 2-2). The model performance was assessed using the R_{adj} and the residual's statistical indices (section 2.3.2d) associated with the measurement's average of the two validation experiments. The obtained results are summarized in Table 2-5 and in Fig. 2-6.

Table 2-5. Model structure Zenteno 10325 averaged performance indices when applied for predictions over the *pilot-scale* experiments.

R_{adj} YAN	R_{adj} Glucose	R_{adj} Fructose	R_{adj} Sugar	R_{adj} Density	Not normal residuals	% Correlated residuals
-3.90 ± 6.12	0.95 $\pm$ 0.01	0.95 $\pm$ 0.01	0.98 $\pm$ 0.02	0.98 $\pm$ 0.02	YAN	0%

The predictive simulations' performance with Zenteno-10325 (Table 2-5) is much better than that of the original model structure (Table 2-3) for the validation experimental pilot-scale dataset. The average R_{adj} associated with the density

increased from 0.79 to 0.98, and its standard deviation was reduced by around 90%; hence, density predictions were more accurate and reliable.

These results are particularly relevant for industrial applications since density is the principal variable used by enologists in decision-making during the alcoholic fermentation, and its use in new control applications for this process has shown relevant interest (Sablayrolles, 2009). Conversely, low YAN predictive performance persisted; Fig. 2-3 and Fig. 2-6 confirm that this is related to the few measurements at the beginning of the fermentations (2-3 measurements before the YAN stagnation point). This limitation also explains a biased estimation of the biomass generation at the beginning of the fermentation, resulting in high prediction errors of the initial sugar consumption. Predictions of sugar-related measurements in later stages of the fermentation were not seriously affected by these initial inaccuracies. Model-based design of experiments can help establish sampling points that would lead to better predictive models.

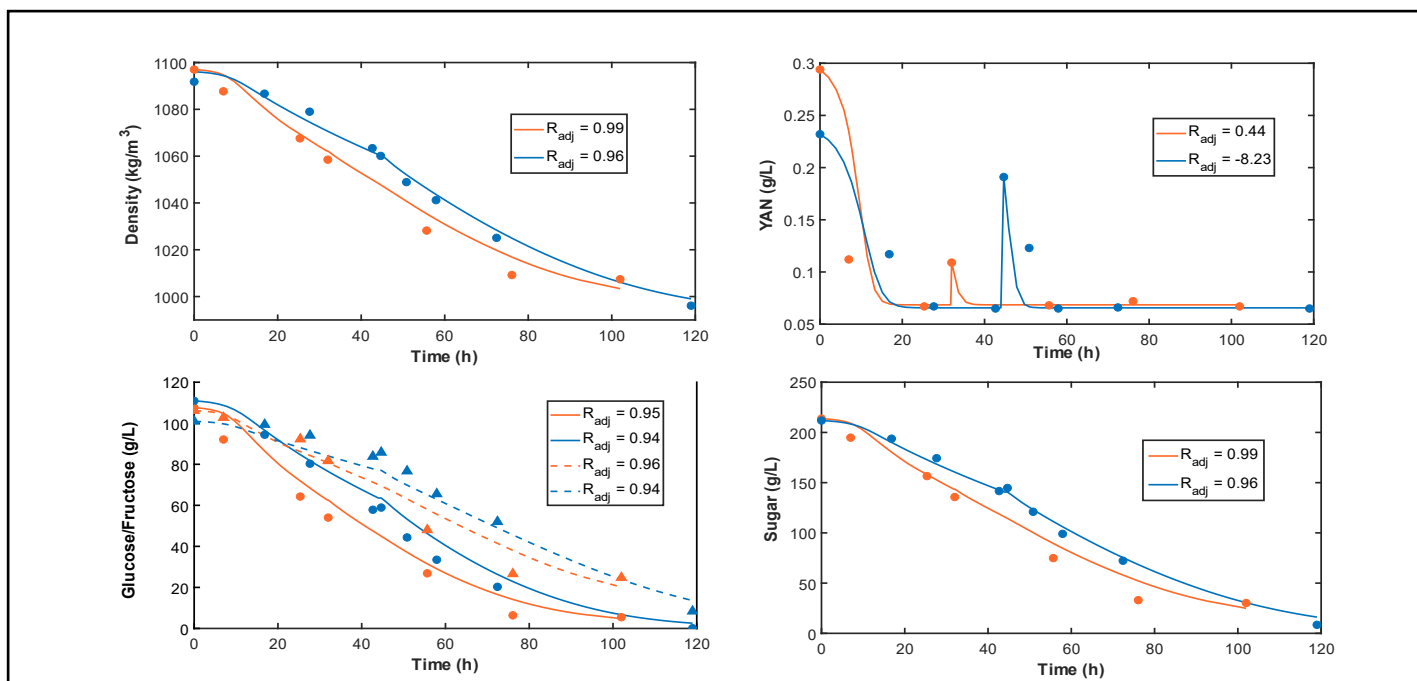


Figure 2-6. Application of model structure Zenteno 10325 over the pilot-scale validation experiments simulation. Blue and orange (line and symbols) represent experiments CII-LO(27) and CII-MA(01), respectively. Symbols represent experimental data, while lines correspond to simulations. In the Glucose/Fructose sub-figure, dashed lines and triangles are associated to Glucose consumption kinetics, while solid lines and circles are associated to Fructose consumption kinetics. The sugar state was simulated as the combination of the Glucose and Fructose model states.

2.5. Conclusion

This paper proposed a robustness-guided workflow for reparametrization assessment and selection of ODE dynamic models to optimize their predictive capability, reliability, and flexibility to fit different experimental data. The procedure consists of four stages and was applied to select a wine fermentation model (starting from a pool of two models) that achieved a good fit using the laboratory-scale and pilot-scale experimental data. During stage I, model structures varying in their free and fixed parameters were characterized through robustness-related indices using the HIPPO algorithm; here, we explored a total of 4092 and 11958 structures corresponding to each evaluated model. Afterward, these model structures were re-processed during stage II, where we selected those models which accomplished minimum robustness requirements, reducing the number of characterized models to a total of 26 and 5 “viable” model structures. When further evaluated using MCDM techniques during stage III, we were able to discard most of the models, leaving us with a small selection of highly robust models. Finally, in stage IV, we used the MCDM-selected models for simulation with the complete laboratory-scale validation and calibration datasets. All the above yielded the selection of model structure “Zenteno-10325” as the best-overall model structures, where only 6 of the initial 14 remained as free parameters for regression purposes. When used for prediction generation over the validation dataset, model structures Zenteno-10325 showed a 5% increase in global predictive performance when compared to the original models. This model structure was validated with our pilot-scale experimental dataset, where we obtained a significant improvement over quality related to sugar predictions; this was observed as an increase in the average fitting performance (R_{adj}) from 0.79 to 0.92. Nevertheless, the prediction of nitrogen-related variables was inaccurate due to insufficient measurements at the beginning of the fermentation. Moreover, we also discuss a method for preliminary structural identifiability assessment, which could be implemented in HIPPO to significantly reduce processing time (four our

example, we estimate this could reduce the number of explored model structures by around 80%). Overall, when a limited data structure is at hand, the proposed workflow yields a reliable minimum model structure with identifiable, independent, and significant fitting parameters. The selected model structure, applied at different scales, accurately predicts the process behavior and achieves good fit of the experimental data. As a final remark, though applied to models lacking features to adequately represent industrial wine fermentation systems, this method states a framework to guide regression procedures for most ODE-based dynamic process models. If properly used during large-scale fermenter model calibration, we expect a significant increase in overall prediction quality and certainty in these systems, directly impacting the viability for their implementation in industrial applications.

2.6. Acknowledgements

Cristóbal Torrealba: Methodology, Investigation, Formal Analysis, Software, Data Curation, Visualization, Writing – Original Draft. **Ricardo Luna:** Supervision, Conceptualization, Project Administration, Writing – Review & Editing, Validation. **José Cuevas-Valenzuela:** Validation, Writing – Review & Editing, Resources, Funding Acquisition. **José R. Pérez-Correa:** Supervision, Conceptualization, Validation, Writing – Review & Editing.

3. FINAL REMARKS AND FUTURE PERSPECTIVES

Though applied to wine AF modeling, it is important to state that the method generated from this work is fully applicable in dynamic models describing processes out of this topic, thus, highlighting the value of its development. Insights from this work aim to become a useful addition for a modeler's repertoire as it is common to see restrictions limiting data in most industrial developments, ultimately forcing modelers to adapt to generate reliable results. Moreover, applications differing from DPE can be approached using this tool, as model reduction tends to shorten computation time while improving accuracy in methods aiming to evaluate parameter interdependence exploration; this is especially useful in those methods where Monte Carlo simulations are central in their procedures. As an example, though based on expert knowledge, Krausch et al. (2019) noticed that fixing analyzed model parameters lead into drastically reduced parameter confidence interval sizes when performing model-based design of experiments (MBDoE), resulting in dramatically lower computation times with higher certainty over obtained results.

Regarding CRI's mission, model structure Zenteno-10325 is expected to be implemented for AF simulation, as results from this work have demonstrated its predictive superiority when compared to models being used in the SmartWinery platform. However, as industrial wine fermentations are the focus of SmartWinery, this model represents just but a stepping-stone pursuing the objective of reliably generating predictions for this process scale. The method leading to Zenteno-10325 has cleared some difficulties related to DPE when restricted data is at hand; situation to be expected in future industrial developments at CRI. This situation encourages the use of the proposed method in models more adequate for industrial winemaking to determine a minimally viable model structure whose parameters are able of being robustly determined, thus, guaranteeing a higher predictive reliability. Additionally, and furtherly aiding the previous, it is expected that the combination of this method with MBDoE procedures is realized at CRI to generate a precise sampling protocol for future vintage

seasons (optimal sampling). Given the above, results from this work represent a highly valuable asset for SmartWinery's further improvement.

REFERENCES

- Akaike, H. (1998). *Information Theory and an Extension of the Maximum Likelihood Principle* (pp. 199–213). Springer, New York, NY. https://doi.org/10.1007/978-1-4612-1694-0_15
- Armstrong, M. (2020). *Cheat sheet: What is a Digital Twin?* IBM Business Operations Blog. <https://www.ibm.com/blogs/internet-of-things/iot-cheat-sheet-digital-twin/>
- Bartsch, J., Borzì, A., Schenk, C., Schmidt, D., Müller, J., Schulz, V., & Velten, K. (2019). An extended model of wine fermentation including aromas and acids. *ArXiv*. <http://arxiv.org/abs/1901.03659>
- Benyahia, B., Lakerveld, R., & Barton, P. I. (2012). A plant-wide dynamic model of a continuous pharmaceutical process. *Industrial and Engineering Chemistry Research*, 51(47), 15393–15412. <https://doi.org/10.1021/ie3006319>
- Bonate, P. L. (2011). The Art of Modeling. In *Pharmacokinetic-Pharmacodynamic Modeling and Simulation* (pp. 1–60). Springer US. https://doi.org/10.1007/978-1-4419-9485-1_1
- Bordons, C., & Núñez-Reyes, A. (2008). Model based predictive control of an olive oil mill. *Journal of Food Engineering*. <https://doi.org/10.1016/j.jfoodeng.2007.04.011>
- Boulton, R. (1980). The Prediction of Fermentation Behavior by a Kinetic Model. *American Journal of Enology and Viticulture*, 31(1).
- Brossard, N., Cai, H., Osorio, F., Bordeu, E., & Chen, J. (2016). “Oral” Tribological Study on the Astringency Sensation of Red Wines. *Journal of Texture Studies*, 47(5), 392–402. <https://doi.org/10.1111/jtxs.12184>
- Brown, S. W., Oliver, S. G., Harrison, D. E. F., & Righelato, R. C. (1981). Ethanol inhibition of yeast growth and fermentation: Differences in the magnitude and complexity of the effect. *European Journal of Applied Microbiology and Biotechnology*, 11(3), 151–155. <https://doi.org/10.1007/BF00511253>
- Cerda-Drago, T. G., Agosin, E., & Pérez-Correa, J. R. (2016). Modelling the oxygen dissolution rate during oenological fermentation. *Biochemical Engineering Journal*, 106, 97–106. <https://doi.org/10.1016/j.bej.2015.10.014>
- Childs, B. C., Bohlscheid, J. C., & Edwards, C. G. (2015). Impact of available nitrogen and sugar concentration in musts on alcoholic fermentation and subsequent wine spoilage by *Brettanomyces bruxellensis*. *Food Microbiology*. <https://doi.org/10.1016/j.fm.2014.10.006>

- Coleman, M. C., Fish, R., & Block, D. E. (2007). Temperature-dependent kinetic model for nitrogen-limited wine fermentations. *Applied and Environmental Microbiology*. <https://doi.org/10.1128/AEM.00670-07>
- Cramer, A. C., Vlassides, S., & Block, D. E. (2002). Kinetic model for nitrogen-limited wine fermentations. *Biotechnology and Bioengineering*, 77(1), 49–60. <https://doi.org/10.1002/bit.10133>
- Detle, H., Melas, V. B., Pepelyshev, A., & Strigul, N. (2005). Robust and efficient design of experiments for the Monod model. *Journal of Theoretical Biology*, 234(4), 537–550. <https://doi.org/10.1016/j.jtbi.2004.12.011>
- Dinnella, C., Recchia, A., Tuorila, H., & Monteleone, E. (2011). Individual astringency responsiveness affects the acceptance of phenol-rich foods. *Appetite*, 56(3), 633–642. <https://doi.org/10.1016/j.appet.2011.02.017>
- Egea, J. A., Balsa-Canto, E., García, M. S. G., & Banga, J. R. (2009). Dynamic optimization of nonlinear processes with an enhanced scatter search method. *Industrial and Engineering Chemistry Research*, 48(9), 4388–4401. <https://doi.org/10.1021/ie801717t>
- Farebrother, R. W. (1980). The Durbin-Watson Test for Serial Correlation when there is no Intercept in the Regression. *Econometrica*. <https://doi.org/10.2307/1912825>
- Fraga, H. (2020). Climate Change: A New Challenge for the Winemaking Sector. *Agronomy*, 10(10), 1465. <https://doi.org/10.3390/agronomy10101465>
- Franceschini, G., & Macchietto, S. (2008). Model-based design of experiments for parameter precision: State of the art. *Chemical Engineering Science*, 63(19), 4846–4872. <https://doi.org/10.1016/j.ces.2007.11.034>
- Galitsky, C., Radspieler, A., Worrell, E., Healy, P., & Zechiel, S. (2005). Benchmarking and self-assessment in the wine industry. *Proceedings ACEEE Summer Study on Energy Efficiency in Industry*, 36–47.
- Gurbey, A. P. (2020). Climate Change Problems in Agricultural Landscape Areas: Eastern Thrace Vineyards. *Journal of Environmental Protection and Ecology*, 21(3), 1090–1097.
- Hao, H., Zak, D. E., Sauter, T., Schwaber, J., & Ogunnaike, B. A. (2006). Modeling the VPAC2-activated cAMP/PKA signaling pathway: From receptor to circadian clock gene induction. *Biophysical Journal*. <https://doi.org/10.1529/biophysj.105.065250>
- Henriques, D., Alonso-del-Real, J., Querol, A., & Balsa-Canto, E. (2018).

Saccharomyces cerevisiae and *S. kudriavzevii* Synthetic Wine Fermentation Performance Dissected by Predictive Modeling. *Frontiers in Microbiology*, 9(FEB), 88. <https://doi.org/10.3389/fmicb.2018.00088>

Holt, H. E., Francis, I. L., Field, J., Herderich, M. J., & Iland, P. G. (2008). Relationships between berry size, berry phenolic composition and wine quality scores for Cabernet Sauvignon (*Vitis vinifera* L.) from different pruning treatments and different vintages. *Australian Journal of Grape and Wine Research*, 14(3), 191–202. <https://doi.org/10.1111/j.1755-0238.2008.00019.x>

Holzberg, I., Finn, R. K., & Steinkraus, K. H. (1967). A kinetic study of the alcoholic fermentation of grape juice. *Biotechnology and Bioengineering*, 9(3), 413–427. <https://doi.org/10.1002/bit.260090312>

Jaqaman, K., & Danuser, G. (2006). Linking data to models: Data regression. In *Nature Reviews Molecular Cell Biology* (Vol. 7, Issue 11, pp. 813–819). Nature Publishing Group. <https://doi.org/10.1038/nrm2030>

Jones, D. F., Mirrazavi, S. K., & Tamiz, M. (2002). Multi-objective meta-heuristics: An overview of the current state-of-the-art. *European Journal of Operational Research*, 137(1), 1–9. [https://doi.org/10.1016/S0377-2217\(01\)00123-0](https://doi.org/10.1016/S0377-2217(01)00123-0)

Koch, K. (2013). *Parameter estimation and hypothesis testing in linear models*. <https://books.google.com/books?hl=es&lr=&id=n3bvCAAQBAJ&oi=fnd&pg=PA2&dq=Parameter+Estimation+and+Hypothesis+Testing+in+Linear+Model&ots=JREPI7Kv7S&sig=gK9buJBQikyCKBLDhTvf52dctwk>

Krausch, N., Barz, T., Sawatzki, A., Gruber, M., Kamel, S., Neubauer, P., & Cruz Bournazou, M. N. (2019). Monte Carlo Simulations for the Analysis of Non-linear Parameter Confidence Intervals in Optimal Experimental Design. *Frontiers in Bioengineering and Biotechnology*, 7, 122. <https://doi.org/10.3389/fbioe.2019.00122>

Kreutz, C., Raue, A., Kaschek, D., & Timmer, J. (2013). Profile likelihood in systems biology. *FEBS Journal*, 280(11), 2564–2571. <https://doi.org/10.1111/febs.12276>

Lehmann, E., & Romano, J. (2006). *Testing statistical hypotheses*. https://books.google.com/books?hl=es&lr=&id=K6t5qn-SEp8C&oi=fnd&pg=PR7&dq=Lehmann+%26+Romano+2005+Testing+Statistical+Hypothesis&ots=dBosfnKdi8&sig=Nyf2x0v9HHkysW9DfD2R24a4w_I

Luna, R., Araya, M., Caris J., & Cuevas-Valenzuela, J. (2020). A Digital Platform for the Management of Grapes and Wine Quality in the Winery. In *2020 39th International Conference of the Chilean Computer Science Society (SCCC)*(pp- 1-7). IEEE.

Ma, W., Guo, A., Zhang, Y., Wang, H., Liu, Y., & Li, H. (2014). A review on astringency and bitterness perception of tannins in wine. In *Trends in Food Science and*

Technology (Vol. 40, Issue 1, pp. 6–19). Elsevier Ltd.
<https://doi.org/10.1016/j.tifs.2014.08.001>

Mahdi, Y., Mouheb, A., & Oufer, L. (2009). A dynamic model for milk fouling in a plate heat exchanger. *Applied Mathematical Modelling*, 33(2), 648–662.
<https://doi.org/10.1016/j.apm.2007.11.030>

Miao, H., Xia, X., Perelson, A. S., & Wu, H. (2011). On identifiability of nonlinear ODE models and applications in viral dynamics. In *SIAM Review* (Vol. 53, Issue 1, pp. 3–39). Society for Industrial and Applied Mathematics.
<https://doi.org/10.1137/090757009>

Miller, K. V., & Block, D. E. (2020). A review of wine fermentation process modeling. In *Journal of Food Engineering* (Vol. 273, p. 109783). Elsevier Ltd.
<https://doi.org/10.1016/j.jfoodeng.2019.109783>

Miller, K. V., Noguera, R., Beaver, J., Oberholster, A., & Block, D. E. (2020). A combined phenolic extraction and fermentation reactor engineering model for multiphase red wine fermentation. *Biotechnology and Bioengineering*, 117(1), 109–116.
<https://doi.org/10.1002/bit.27178>

Miller, K. V., Oberholster, A., & Block, D. E. (2019). Creation and validation of a reactor engineering model for multiphase red wine fermentations. *Biotechnology and Bioengineering*, 116(4), 781–792. <https://doi.org/10.1002/bit.26874>

Mira de Orduña, R. (2010). Climate change associated effects on grape and wine quality and production. *Food Research International*.
<https://doi.org/10.1016/j.foodres.2010.05.001>

Monod, J. (1949). The Growth of Bacterial Cultures. *Annual Review of Microbiology*, 3(1), 371–394. <https://doi.org/10.1146/annurev.mi.03.100149.002103>

Pantelides, C. C., & Renfro, J. G. (2013). The online use of first-principles models in process operations: Review, current status and future needs. *Computers and Chemical Engineering*, 51, 136–148. <https://doi.org/10.1016/j.compchemeng.2012.07.008>

Pizarro, F., Vargas, F. A., & Agosin, E. (2007). A systems biology perspective of wine fermentations. *Yeast*, 24(11), 977–991. <https://doi.org/10.1002/yea.1545>

Qin, S. J., & Badgwell, T. A. (2003). A survey of industrial model predictive control technology. *Control Engineering Practice*. [https://doi.org/10.1016/S0967-0661\(02\)00186-7](https://doi.org/10.1016/S0967-0661(02)00186-7)

Ribéreau-Gayon, P., Glories, Y., Maujean, A., & Dubourdieu, D. (2006). Handbook of

Enology. In *Handbook of Enology*. <https://doi.org/10.1002/0470010398>

Rodríguez-Fernández, M., Balsa-Canto, E., Egea, J. A., & Banga, J. R. (2007). Identifiability and robust parameter estimation in food process modeling: Application to a drying model. *Journal of Food Engineering*, 83(3), 374–383. <https://doi.org/10.1016/j.jfoodeng.2007.03.023>

Saa, P. A., Moenne, M. I., Pérez-Correa, J. R., & Agosin, E. (2012). Modeling oxygen dissolution and biological uptake during pulse oxygen additions in oenological fermentations. *Bioprocess and Biosystems Engineering*, 35(7), 1167–1178. <https://doi.org/10.1007/s00449-012-0703-7>

Saa, P. A., & Nielsen, L. K. (2017). Formulation, construction and analysis of kinetic models of metabolism: A review of modelling frameworks. In *Biotechnology Advances* (Vol. 35, Issue 8, pp. 981–1003). Elsevier Inc. <https://doi.org/10.1016/j.biotechadv.2017.09.005>

Sablairolles, J. M. (2009). Control of alcoholic fermentation in winemaking: Current situation and prospect. *Food Research International*, 42(4), 418–424. <https://doi.org/10.1016/j.foodres.2008.12.016>

Sacher, J., Saa, P., Cárcamo, M., López, J., Gelmi, C. A., & Pérez-Correa, R. (2011). Improved calibration of a solid substrate fermentation model. *Electronic Journal of Biotechnology*, 14(5). <https://doi.org/10.2225/vol14-issue5-fulltext-7>

Sainz, J., Pizarro, F., Pérez-Correa, J. R., & Agosin, E. (2003). Modeling of yeast metabolism and process dynamics in batch fermentation. *Biotechnology and Bioengineering*, 81(7), 818–828. <https://doi.org/10.1002/bit.10535>

Saitua, F., Torres, P., Pérez-Correa, J. R., & Agosin, E. (2017). Dynamic genome-scale metabolic modeling of the yeast *Pichia pastoris*. *BMC Systems Biology*, 11(1), 27. <https://doi.org/10.1186/s12918-017-0408-2>

Sánchez, B. J., Pérez-Correa, J. R., & Agosin, E. (2014a). Construction of robust dynamic genome-scale metabolic model structures of *Saccharomyces cerevisiae* through iterative re-parameterization. *Metabolic Engineering*, 25, 159–173. <https://doi.org/10.1016/j.ymben.2014.07.004>

Sánchez, B. J., Soto, D. C., Jorquera, H., Gelmi, C. A., & Pérez-Correa, J. R. (2014b). HIPPO: An iterative reparametrization method for identification and calibration of dynamic bioreactor models of complex processes. *Industrial and Engineering Chemistry Research*, 53(48), 18514–18525. <https://doi.org/10.1021/ie501298b>

Sánchez, B. J., Zhang, C., Nilsson, A., Lahtvee, P., Kerkhoven, E. J., & Nielsen, J.

(2017). Improving the phenotype predictions of a yeast genome-scale metabolic model by incorporating enzymatic constraints. *Molecular Systems Biology*, 13(8), 935. <https://doi.org/10.15252/msb.20167411>

Setford, P. C., Jeffery, D. W., Grbin, P. R., & Muhlack, R. A. (2017). Factors affecting extraction and evolution of phenolic compounds during red wine maceration and the role of process modelling. *Trends in Food Science & Technology*, 69, 106-117.

Setford, P. C., Jeffery, D. W., Grbin, P. R., & Muhlack, R. A. (2019). Mathematical modelling of anthocyanin mass transfer to predict extraction in simulated red wine fermentation scenarios. *Food Research International*, 121, 705-713.

Schmid, F., Schadt, J., Jiranek, V., & Block, D. E. (2009). Formation of temperature gradients in large- and small-scale red wine fermentations during cap management. *Australian Journal of Grape and Wine Research*, 15(3), 249–255. <https://doi.org/10.1111/j.1755-0238.2009.00053.x>

Schwinn, M., Durner, D., Delgado, A., & Fischer, U. (2019). Distribution of yeast cells, temperature, and fermentation by-products in white wine fermentations. *American Journal of Enology and Viticulture*, 70(4), 339–350. <https://doi.org/10.5344/ajev.2019.18092>

Serebrinsky, K., Hirmas, B., Munizaga, J., & Pedreros, F. (2019). Model structures for batch and fed-batch ethanol fermentations. *IEEE CHILEAN Conference on Electrical, Electronics Engineering, Information and Communication Technologies, CHILECON 2019*. <https://doi.org/10.1109/CHILECON47746.2019.8988018>

Singleton, V. L., & Trousdale, E. K. (1992). Anthocyanin-Tannin Interactions Explaining Differences in Polymeric Phenols Between White and Red Wines. *American Journal of Enology and Viticulture*, 43(1).

Stigter, J. D., & Molenaar, J. (2015). A fast algorithm to assess local structural identifiability. *Automatica*, 58, 118–124. <https://doi.org/10.1016/j.automatica.2015.05.004>

Streif, S., Kim, K. K. K., Rumschinski, P., Kishida, M., Shen, D. E., Findeisen, R., & Braatz, R. D. (2013). Robustness analysis, prediction and estimation for uncertain biochemical networks. *IFAC Proceedings Volumes (IFAC-PapersOnline)*. <https://doi.org/10.3182/20131218-3-IN-2045.00190>

Suárez-Lepe, J. A., & Morata, A. (2012). New trends in yeast selection for winemaking. In *Trends in Food Science and Technology* (Vol. 23, Issue 1, pp. 39–50). Elsevier. <https://doi.org/10.1016/j.tifs.2011.08.005>

Tao, F., Zhang, H., Liu, A., & Nee, A. Y. C. (2019). Digital Twin in Industry: State-of-the-Art. *IEEE Transactions on Industrial Informatics*, 15(4), 2405–2415.

<https://doi.org/10.1109/TII.2018.2873186>

Unterkofler, J., Muhlack, R. A., & Jeffery, D. W. (2020). Processes and purposes of extraction of grape components during winemaking: current state and perspectives. In *Applied Microbiology and Biotechnology* (Vol. 104, Issue 11, pp. 4737–4755). Springer. <https://doi.org/10.1007/s00253-020-10558-3>

Vargas, F. A., Pizarro, F., Pérez-Correa, J. R., & Agosin, E. (2011). Expanding a dynamic flux balance model of yeast fermentation to genome-scale. *BMC Systems Biology*, 5(1), 75. <https://doi.org/10.1186/1752-0509-5-75>

Vlassides, S., & Block, D. E. (2000). Evaluation of Cell Concentration Profiles and Mixing in Unagitated Wine Fermentors. *American Journal of Enology and Viticulture*, 51(1).

Wang, Z., & Rangaiah, G. P. (2017b). Application and Analysis of Methods for Selecting an Optimal Solution from the Pareto-Optimal Front obtained by Multiobjective Optimization. *Industrial and Engineering Chemistry Research*, 56(2), 560–574. <https://doi.org/10.1021/acs.iecr.6b03453>

Webb, L. B., Whetton, P. H., & Barlow, E. W. R. (2007). Modelled impact of future climate change on the phenology of winegrapes in Australia. *Australian Journal of Grape and Wine Research*, 13(3), 165–175. <https://doi.org/10.1111/j.1755-0238.2007.tb00247.x>

Yacco, R. S., Watrelot, A. A., & Kennedy, J. A. (2016). Red Wine Tannin Structure-Activity Relationships during Fermentation and Maceration. *Journal of Agricultural and Food Chemistry*, 64(4), 860–869. <https://doi.org/10.1021/acs.jafc.5b05058>

Zenteno, M. I., Pérez-Correa, J. R., Gelmi, C. A., & Agosin, E. (2010). Modeling temperature gradients in wine fermentation tanks. *Journal of Food Engineering*, 99(1), 40–48. <https://doi.org/10.1016/j.jfoodeng.2010.01.033>

Zhang, Q., Wu, D., Lin, Y., Wang, X., Kong, H., & Tanaka, S. (2015). Substrate and product inhibition on yeast performance in ethanol fermentation. *Energy and Fuels*, 29(2), 1019–1027. <https://doi.org/10.1021/ef502349v>

APPENDIX

APPENDIX A: REDUCED ZENTENO MODEL

The reduced Zenteno model was obtained from Zenteno et al. (2010), with temperature incorporated as an input for the model. No compartmentalization was used, since in both laboratory and pilot-scale experiments, temperature gradients were considered not to be significant. CO₂ generation was not incorporated into the model, as it did not affect other state variables. Density was estimated using sugar levels (glucose + fructose) and a linear interpolation obtained by adjustment using previous experiments in our laboratory. Therefore, the reduced Zenteno model is constructed as follows:

Biomass generation and decay:

$$\frac{dX}{dt} = (\mu - k_d) \cdot X \quad (\text{A.1})$$

Nitrogen consumption:

$$\frac{dN}{dt} = -\frac{\mu}{Y_{XN}} \cdot X \quad (\text{A.2})$$

Glucose consumption:

$$\frac{dG}{dt} = -\left(\frac{\mu}{Y_{XG}} + \frac{\beta_G}{Y_{EG}} + m \cdot \frac{G}{G+F}\right) \cdot X \quad (\text{A.3})$$

Fructose consumption:

$$\frac{dF}{dt} = -\left(\frac{\mu}{Y_{XF}} + \frac{\beta_F}{Y_{EF}} + m \cdot \frac{F}{G+F}\right) \cdot X \quad (\text{A.4})$$

Ethanol production:

$$\frac{dE}{dt} = (\beta_G + \beta_F) \cdot X \quad (\text{A.5})$$

With:

Specific growth rate:

$$\mu = \mu_{max} \cdot \frac{N}{N + K_N(T)} \quad (\text{A.6})$$

Maximum growth rate:

$$\mu_{max} = \mu_0 \cdot \exp\left(\frac{E_{ac} \cdot (T - 300)}{300 \cdot R \cdot T}\right) \quad (\text{A.7})$$

Biomass decay rate (when $T > T_D$):

$$\mu_{max} = k_{d0} \cdot \exp\left(C_{de} \cdot E + E_{td} \cdot \left(\frac{T - 305.65}{306.65 \cdot R \cdot T}\right)\right) \quad (\text{A.8})$$

Thermal death ethanol-related threshold:

$$T_D = -10^{-4} \cdot E^3 + 0.0049 \cdot E^2 - 0.1279 \cdot E + 315.89 \quad (\text{A.9})$$

Ethanol production rate from fructose:

$$\beta_F = \beta_{Fmax} \cdot \frac{F}{F + K_F(T)} \cdot \frac{K_{IG}(T)}{G + K_{IG}(T)} \cdot \frac{K_{IE}(T)}{E + K_{IE}(T)} \quad (\text{A.10})$$

Ethanol production rate from glucose

$$\beta_G = \beta_{Gmax} \cdot \frac{G}{G + K_G(T)} \cdot \frac{K_{IE}(T)}{R + K_{IE}(T)} \quad (\text{A.11})$$

Specific cell maintenance rate:

$$m = m_0 \cdot \exp\left(\frac{E_{am} \cdot (T - 293.3)}{293.3 \cdot R \cdot T}\right) \quad (\text{A.12})$$

Maximum ethanol production rates:

$$\beta_{i,max} = \beta_{i0} \cdot \exp\left(\frac{E_{afe} \cdot (T - 296.15)}{296.15 \cdot R \cdot T}\right), \quad i = G, F \quad (\text{A.13})$$

Temperature inhibition parameters:

$$K_i(T) = K_{i0} \cdot \exp\left(\frac{E_{aki} \cdot (T - T^*)}{293.3 \cdot R \cdot T}\right), \quad i = N, G, F, iG, iE \quad (\text{A.14})$$

For all parameters values not included in the parameter estimation, i.e., not in the table from Appendix B, refer to Zenteno et al. (2010).

APPENDIX B: MODEL STRUCTURE ZENTENO-10325

As described above, model structures are models where a specific set of estimation parameters are fixed given their lack of robustness for model calibration. Model structure Zenteno 10325 corresponds to the model depicted in Appendix A; fixed, and estimated parameters are shown in the following table (Table B-1).

Table B-1. Parameter enumeration corresponding to the reduced Zenteno model. Optimal values for fixed and free parameters of Zenteno model structure 10325 are also indicated.

Parameter N°	Model parameter	Fixed value	Estimated value
θ_1	μ_0	-	0.54 h ⁻¹
θ_2	β_{G0}	-	0.30 h ⁻¹
θ_3	β_{F0}	-	0.26 h ⁻¹
θ_4	K_{N0}	-	010 kg m ⁻³
θ_5	K_{G0}	8.84 kg m ⁻³	-
θ_6	K_{F0}	-	11.97 kg m ⁻³
θ_7	K_{iG0}	-	56.65 kg m ⁻³
θ_8	K_{iE0}	36.80 kg m ⁻³	-
θ_9	K_{d0}	3.94 · 10 ⁻⁵ h ⁻¹	-
θ_{10}	Y_{XN}	21.32	-
θ_{11}	Y_{XG}	1.66	-
θ_{12}	Y_{XF}	1.41	-
θ_{13}	Y_{EG}	0.58	-
θ_{14}	Y_{EF}	0.63	-

APPENDIX C: MCDM RESULTS WHEN APPLIED TO GROUP I STRUCTURES

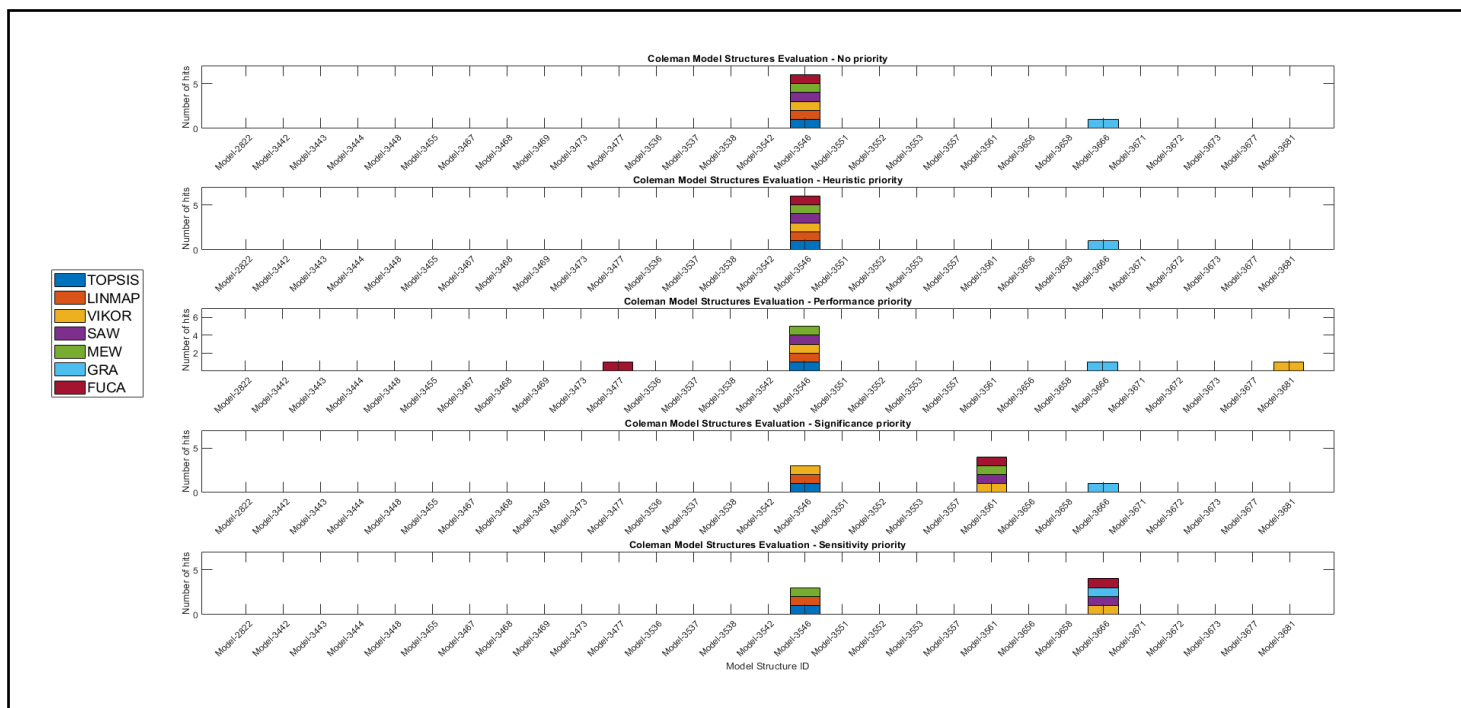


Figure C-1. MCDM results when applied to model structures from group I (Coleman originated models). Stacked bars represent votes committed by each of the MCDM methods listed in the figure legend (Table 2-1). Overall, model structures 3545, 3561 and 3666 represent good levels of robustness, as these models were the most voted in one or more scenario

APPENDIX D: COLEMAN MODEL STRUCTURAL IDENTIFIABILITY ANALYSIS RESULT

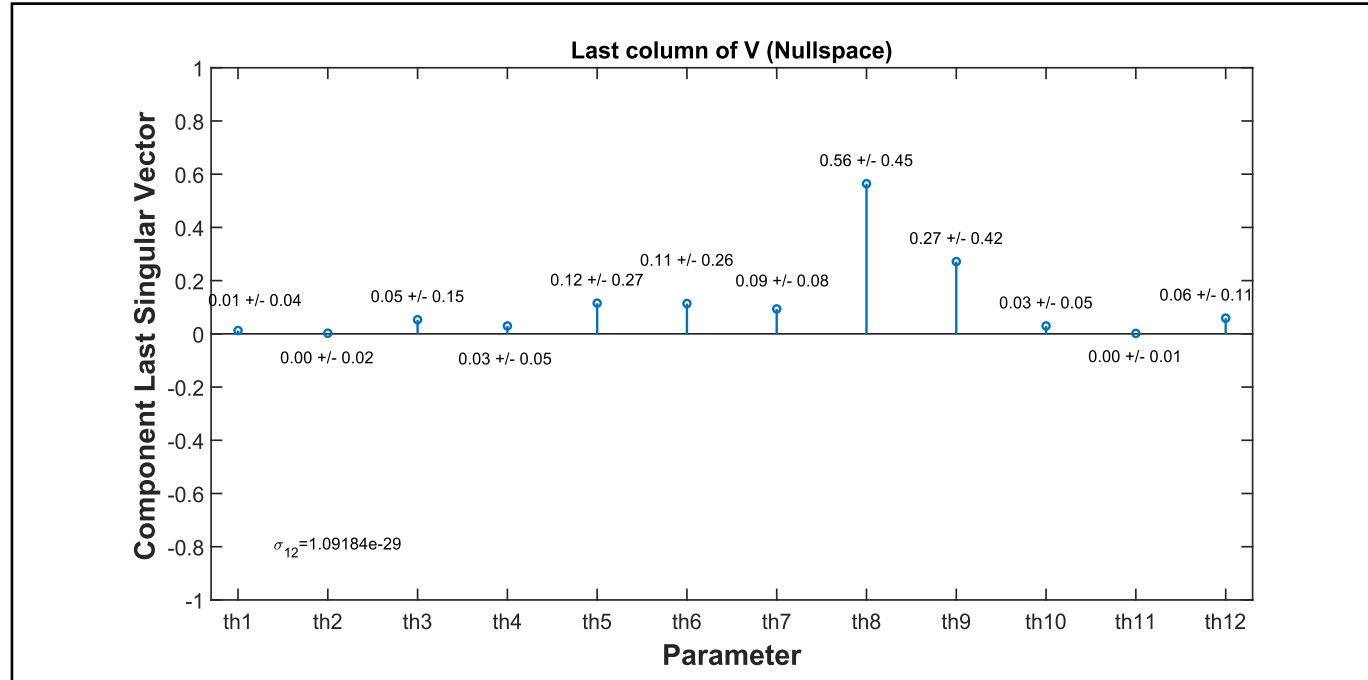


Figure D-1. Mean absolute values of the components of the last singular vector of the ROSM obtained from applying the method proposed by Stigter & Molenaar (2015) over the Coleman model. A total of 500 iterations were performed, varying nominal parameter values. Values distant from zero represent problematic parameters given the data structure used for calibration. σ_{14} corresponds to the singular value associated to the last column in the SVD decomposition of the ROSM.

