



**A Method of Surrogate Model Construction which Leverages
Lower-Fidelity Information using Space Mapping Techniques**

THESIS

Jason W. Thomas, Capt, USAF

AFIT-ENY-14-M-46

**DEPARTMENT OF THE AIR FORCE
AIR UNIVERSITY**

AIR FORCE INSTITUTE OF TECHNOLOGY

Wright-Patterson Air Force Base, Ohio

DISTRIBUTION A: APPROVED FOR PUBLIC RELEASE; DISTRIBUTION IS UNLIMITED

The views expressed in this thesis are those of the author and do not reflect the official policy or position of the United States Air Force, Department of Defense, or the United States Government. This material is declared a work of the U.S. Government and is not subject to copyright protection in the United States.

A Method of Surrogate Model Construction which Leverages Lower-Fidelity
Information using Space Mapping Techniques

THESIS

Presented to the Faculty
Department of Aeronautics and Astronautics
Graduate School of Engineering and Management
Air Force Institute of Technology
Air University
Air Education and Training Command
In Partial Fulfillment of the Requirements for the
Degree of Master of Science in Aeronautical Engineering

Jason W. Thomas, B.S.
Capt, USAF

March 2014

A Method of Surrogate Model Construction which Leverages Lower-Fidelity Information using
Space Mapping Techniques

Jason W. Thomas, BS
Captain, USAF

Approved:

//signed// Lt Col Jeremy Agte, PhD (Chair)	12 March 2014 Date
//signed// James W. Chrissis, PhD (Member)	12 March 2014 Date
//signed// Lt Col Anthony Deluca, PhD (Member)	12 March 2014 Date

Abstract

A new method of surrogate construction is developed and applied to a pair of computational tools used in the field of aircraft design. This new method involves the pairing of data sampled from the analytical model of interest with the execution of a similar analysis performed at a lower level of fidelity. This pairing is accomplished through the use of a space mapping technique, which is a process where the design space of a lower fidelity model is aligned a higher fidelity model. The intent of applying space mapping techniques to the field of surrogate construction is to leverage the information about a system's performance present at a lower fidelity level to bolster the predictive accuracy of a surrogate model based upon sampled data at a higher fidelity level. The results from the pairing of computational tools used in this research show modest gains in predictive accuracy for many of the cases investigated when compared to existing surrogate methodologies.

Acknowledgments

I would like to thank my wife for the support she's shown me over the course of my graduate studies. She is my source for inspiration and the only reason I don't wear sweatpants in public.

I would also like to thank Lt Col Agte, Dr. Alyanak, and Dr. Camberos for their guidance and support of this research. I would have been hopelessly lost had I not been able to tap into their collective knowledge and experience in the fields of optimization and aircraft design.

Jason W. Thomas

Table of Contents

	Page
Abstract	iv
Acknowledgments	v
I Motivation	1
II Background and Theory	5
Least-Squares Projections.....	5
Polynomial Response Methodology	7
Kriging	10
Space Mapping.....	13
III Methodology	16
Assumptions and Limitations.....	16
Modified Space Mapping Approach	18
Modified Space Mapping Algorithm	20
Surrogate Construction Through Space Mapping.....	29
IV Conceptual Applications	31
High-fidelity and Low-fidelity Model Pairs	31
Surrogate Model Comparison.....	34
V Space Mapping Application with ESAV Tools	54
ESAV Space Mapping Implementation	55
ESAV Space Mapping Results	61
VI Conclusions	73
Future Work	74

Appendix	78
References	96

List of Figures

Figure		Page
1	Infographic comparing the cost of execution and the levels of fidelity for different phases in the design cycle of an aircraft	2
2	Generic representation of a surrogate model in parallel with the analysis being approximated	3
3	Illustration of Least-Squares Approximation	7
4	Illustration of a space mapped surrogate model receiving the same inputs as the high-fidelity analysis to produce an approximation of the high-fidelity response.....	19
5	Modified space mapping algorithm, shown with inputs and outputs for each task	21
6	Steps that compose the first task in the modified space mapping process	22
7	Steps that compose the second task in the modified space mapping process	23
8	Steps that compose the final task in the modified space mapping process.....	25
9	Steps specific to the least-squares polynomial form for $\mathbf{P}$	26
10	Steps specific to the nonlinear polynomial form for $\mathbf{P}$	27
11	Steps specific to the nonlinear kriging space mapping form.....	29
12	Difference in lift predictions by model (variables that are held constant are listed in Table 13 in the Appendix)	32
13	Visual comparison of the <code>peaks()</code> function versus the truncated <code>peaks()</code> function	34
14	Visual comparison between two surrogate models with $q = 8$ where (a) is considered a “good” fit while (b) is considered a “poor” fit	36
15	First-order polynomial space-mapped surrogate model using 12 sampled points	37
16	Visual comparison between surrogate models using a first-order and second-order least-squares space mapping	38
17	Total error calculations for each degree of polynomial fit to the space mapping data, $q = 120$	39
18	Average Errors per sampled point for a range of q showing a general decrease in error per point as the polynomial degree is increased.....	39

19	Comparison between the surrogate models constructed assuming (a) a least-squares polynomial and (b) a nonlinear space mapping for $q = 12$	40
20	Comparison between the surrogate models constructed assuming (a) a least-squares polynomial and (b) a nonlinear space mapping for $q = 24$	41
21	Comparison between the surrogate models constructed assuming (a) a least-squares polynomial and (b) a nonlinear space mapping using kriging models for $q = 38$	42
22	Comparison between the surrogate models constructed through (a) a kriging implementation of space mapping and (b) a traditional kriging model acting as a surrogate ..	43
23	High-fidelity model response over a range of values for $\bar{x}_{H_6}$ and $\bar{x}_{H_7}$	44
24	Surrogate model constructed using 12 sampled high-fidelity data points and a linear least-squares space mapping	45
25	Surrogate model constructed using 24 sampled high-fidelity data points and a second-order least-squares space mapping	45
26	Average Errors per sampled point for a range of q 's showing a steep decrease from order 1 to 2, and marginal decreases for subsequent polynomial degrees	46
27	Surrogate models constructed using 36 sampled high-fidelity data points and assuming a (a) second-order and a (b) third-order least-squares space mapping	47
28	Surrogate models constructed using 12 sampled high-fidelity data points and assuming a (a) first-order least-squares and a (b) nonlinear polynomial space mapping	48
29	Comparison of average error per sampled point for a range of q using a first and second-order least-squares polynomial and a nonlinear polynomial space mapping approach ..	48
30	Surrogate models constructed using 38 sampled high-fidelity data points and assuming a (a) kriging implementation of space mapping and a (b) traditional kriging model acting as a surrogate	49
31	Illustration of the subspace, encompassing the peak in the high-fidelity model that is absent in the low-fidelity model, from which data points are sampled for the space mapping process	50
32	Surrogate model constructed using a fifth-order least-squares space mapping approach ..	51
33	Surrogate model constructed using a nonlinear polynomial space mapping approach ..	52
34	Surrogate model constructed using a nonlinear kriging space mapping approach	53
35	Surrogate model constructed using a traditional kriging response surface	53
36	N^2 diagram depicting the various analysis blocks in the ESAV model within SORCER, taken from Ref [14]	55
37	Final representation of the surrogate model constructed through the implementation of the space mapping algorithm	60

38	Histogram and representative normal distribution curve for the percent errors found using the least-squares space-mapped surrogate model in comparison to the high-fidelity response.....	62
39	Histogram and representative normal distribution curve for the percent errors found using the nonlinear space-mapped surrogate model in comparison to the high-fidelity response	63
40	Histogram and representative normal distribution curve for the percent errors found using the polynomial response surrogate (LS PRM) overlaid on the data from the space-mapped (SM) surrogate	64
41	Scatterplot for both the least-squares and nonlinear space-mapped surrogate responses, with the least-squares PRM surrogate response plotted for comparison	65
42	Percent error comparison between the least-squares space-mapping and the PRM surrogate models derived from samples in the second dataset	66
43	Percent error comparison between the least-squares space-mapping and the PRM surrogate models derived from samples in the third dataset	66
44	Histograms and representative normal distribution curves for the first and second-order least-squares surrogate models based on the first of three second-order datasets.	68
45	Histogram and representative normal distribution curve for the nonlinear surrogate model based on the first of three second-order datasets (second-order LS PRM surrogate data plotted for comparison)	69
46	Histograms and representative normal distribution curves for the kriging-based SM surrogate and the traditional kriging surrogate based on the same 58 sample data points from the first kriging dataset	71
47	RMSE plotted against the number of samples on which each kriging surrogate was based	72
48	Bar chart illustrating the coefficient values of the space mapping form shown in Table 10	75
49	Tenth-order polynomial space mapped surrogate for the Case 1 model pairing	82
50	Surrogate constructed using a linear least-squares space mapping for the 3rd model pairing	82
51	Surrogate constructed using a second-order least-squares space mapping for the 3rd model pairing	83
52	Surrogate constructed using a third-order least-squares space mapping for the 3rd model pairing	83
53	Surrogate model constructed using a fourth-order least-squares space mapping approach	84
54	Surrogate model constructed using a sixth-order least-squares space mapping approach	84

55	Comparison between a kriging-based space-mapped surrogate and the corresponding traditional kriging surface built from the sampling locations shown. RMSE values: (b) 2818.6 (c) 9679.2	85
56	Comparison between a kriging-based space-mapped surrogate and the corresponding traditional kriging surface built from the sampling locations shown. RMSE values: (b) 645.1 (c) 910.6	85
57	Comparison between a kriging-based space-mapped surrogate and the corresponding traditional kriging surface built from the sampling locations shown. RMSE values: (b) 459.0 (c) 2379.6.....	85
58	Scatterplot for both the least-squares and nonlinear space-mapped surrogate responses, with the least-squares PRM surrogate response plotted for comparison (second of three first-order datasets).....	87
59	Scatterplot for both the least-squares and nonlinear space-mapped surrogate responses, with the least-squares PRM surrogate response plotted for comparison (third and final first-order dataset).....	87
60	Histograms and representative normal distribution curves for the first and second-order least-squares surrogate models based on the second of three second-order datasets	88
61	Histogram and representative normal distribution curve for the nonlinear surrogate model based on the second of three second-order datasets (second-order LS PRM surrogate data plotted for comparison)	89
62	Histograms and representative normal distribution curves for the first and second-order least-squares surrogate models based on the third of three second-order datasets	90
63	Histogram and representative normal distribution curve for the nonlinear surrogate model based on the third second-order dataset (second-order LS PRM surrogate data plotted for comparison)	91
64	Scatterplot showing the surrogate model predictions against the true high-fidelity response for the first and second-order LS surrogate models for the first of three datasets.....	92
65	Scatterplot showing the surrogate model predictions against the true high-fidelity response for the second-order LS surrogate models as well as the nonlinear polynomial-based SM surrogate for the first of three datasets	92
66	Scatterplot showing the surrogate model predictions against the true high-fidelity response for the first and second-order LS surrogate models for the second of three datasets.....	93
67	Scatterplot showing the surrogate model predictions against the true high-fidelity response for the second-order LS surrogate models as well as the nonlinear polynomial-based SM surrogate for the second of three datasets	93
68	Scatterplot showing the surrogate model predictions against the true high-fidelity response for the first and second-order LS surrogate models for the third dataset	94

69	Scatterplot showing the surrogate model predictions against the true high-fidelity response for the second-order LS surrogate models as well as the nonlinear polynomial-based SM surrogate for the third dataset.....	94
70	Scatterplot showing the surrogate model predictions against the true high-fidelity response for the kriging-based SM surrogate and the traditional kriging surrogate based on the same 58 sample data points from the first kriging dataset	95

List of Tables

Table	Page
1 Noisy Measurements.....	5
2 List of assumptions and associated limitations.....	17
3 Collection of data from the modified space mapping algorithm	20
4 Variables present in the Case 1 analytical models	32
5 Number of high and low-fidelity analysis calls for the first-order space mappings of the Case 1 model pairing	37
6 High and low-fidelity design vector for the ESAV space mapping application	56
7 Shared-fidelity design vector for the ESAV space mapping application	57
8 Upper and lower boundaries for the sampled datasets used in the ESAV space mapping application	58
9 Space mapping coefficient values for the ESAV application without the modification to the algorithm	59
10 Space mapping coefficient values for the ESAV application with the modification to the algorithm	60
11 Surrogate performance summary comparing the various surrogates constructed from the sampled datasets containing 14 data points.....	67
12 Surrogate performance summary comparing the various surrogates constructed from the sampled datasets containing 28 data points.....	70
13 Variable values held constant in the production of Figure 12.....	81
14 Case 1 shared variable values for surrogate model comparisons	81
15 Case 1 variable boundaries.....	81
16 Case 1 variable boundaries.....	81
17 Inputs held constant in the weight predictor for the ESAV space mapping.....	86

A Method of Surrogate Model Construction which Leverages Lower-Fidelity Information using Space Mapping Techniques

I. Motivation

The field of aircraft design has progressed at an astonishing pace since the Wright brothers first designed and flew their Wright Flyer on the sand dunes of Kittyhawk, North Carolina. Advancements in the understanding of aeronautics, propulsion, material mechanics, controls, electronics, computer science, and many other fields have propelled aircraft design from its humble origins to the technological marvels seen in flight today. The challenge of aircraft design is pushing many of these engineering disciplines to the edge of our understanding and experience. Future aircraft require innovative technologies to push the bounds of performance to new heights. These aircraft designs must take advantage of, or be able to withstand, the various physical phenomena affecting the system. To better design and analyze future aircraft, commercial and government agencies interested in the development of future aircraft have invested resources in the development of computational design frameworks that will allow them to model the system and its environment.

Many of these design frameworks incorporate computational tools which yield similar outputs, but calculate their responses at differing levels of fidelity. Fidelity, in the context of analytical design tools, is the degree to which the physical phenomena relevant to the system are accounted for within the analysis. At the highest levels of fidelity, the design team is modeling as much of the physical environment and its interactions with the system as the program can afford in a resource-constrained environment. Higher fidelity in an analytical tool almost always comes at a higher price, be it in man-hours, computational resources, or the time required for the analysis to complete. For this reason, design frameworks may also incorporate lower fidelity computational tools so performance parameters of the system can be approximated at an affordable fidelity level for the appropriate phase of the design cycle. These multifidelity design and analysis programs allow the design team the option of choosing the fidelity of the analysis based on the particular application and the current phase of the design cycle.

The ability to execute an analysis of the design at a lower fidelity level presents advantages and disadvantages to the design of the system. One such advantage offered by incorporating lower fidelity tools is how relatively inexpensive these tools are compared to their higher fidelity counterparts. In many cases, the cost of executing the lower fidelity tools is such that the design

team can afford to execute design variable sweeps to gather trend data, or even apply an optimization routine to the system's design parameters. The information gleaned from the lower fidelity tools is often the basis for decisions made in the conceptual design phase about the system's top-level configuration. One disadvantage associated with using lower fidelity tools to get this information is the risk of excluding potentially relevant analyses contained in the higher fidelity tools. As a result, innovative system configurations may be excluded from consideration in the early phases of the design cycle due to the ignorance of the lower fidelity tools to the more complicated physical phenomena affecting the system. A simplification of a design cycle for an aircraft is shown in Figure 1, which illustrates the cost and level of fidelity associated with each phase of development.

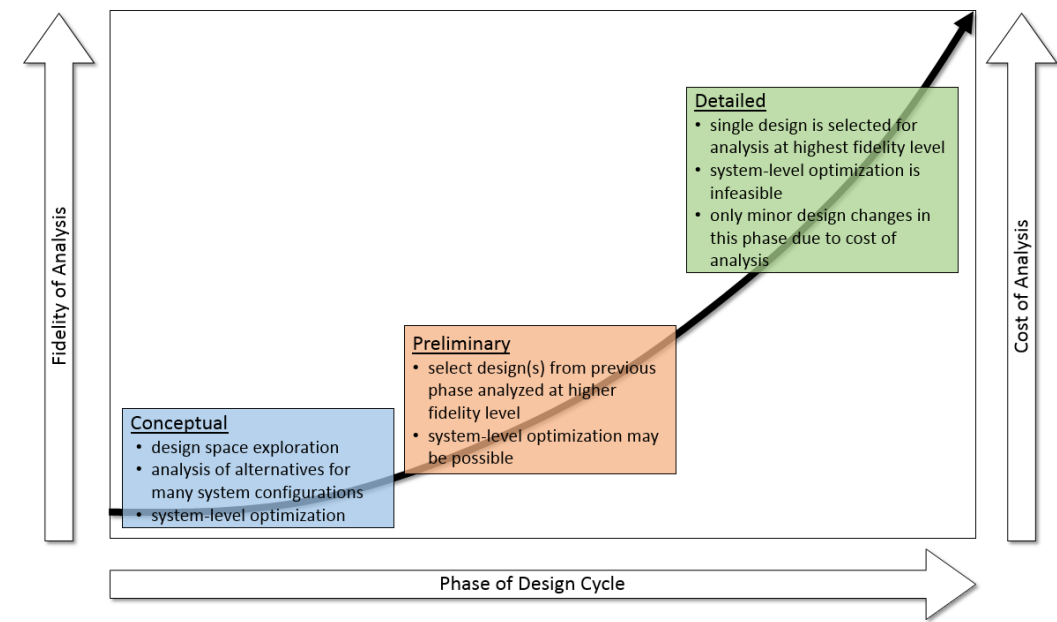


Figure 1. Infographic comparing the cost of execution and the levels of fidelity for different phases in the design cycle of an aircraft

In order to mitigate this risk, the Aerospace Vehicles Directorate of the Air Force Research Laboratory (AFRL/RQ) is searching for ways to represent the information in the higher fidelity levels earlier in the design cycle without incurring the full cost of exhaustively executing the higher fidelity tools. Such a development would allow the tools at the higher levels of fidelity to influence the design space of the system in phases of development where many of the system-level configuration decisions are made. The major obstacle in the development of such a methodology is

approximating the information at the higher levels of fidelity at costs which do not prohibit activities like design space exploration or system-level optimization.

One possible route to bring high-fidelity information into earlier phases of the design cycle is the construction of surrogate models to approximate the high-fidelity response. As is often the case in the conceptual design phase, the actual performance numbers output by the analyses are less important than the trends generated from parametric sweeps of the design parameters. A good surrogate model of a high-fidelity analysis, while only providing approximations at point locations, would retain the same overall trend information as the high-fidelity model. This higher fidelity information would in turn influence the choices made by the design team in an earlier phase of development. A generic representation of a surrogate model is shown in Figure 2, which illustrates how a surrogate model accepts the same inputs as the analytical model being imitated and produces an approximation of the analysis' response.

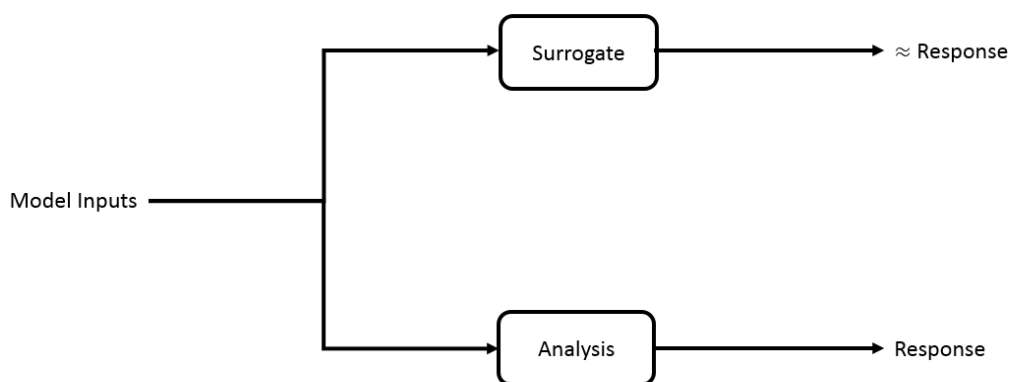


Figure 2. Generic representation of a surrogate model in parallel with the analysis being approximated

Surrogate construction is a field of study in itself, and [1] is an excellent resource on modern techniques. The general process used in surrogate construction involves some method of sampling the high-fidelity model to generate the data on which the surrogate is based. The accuracy of any given surrogate model typically improves as the number of sampled points increases. However, this can be costly especially if the analyses are computationally expensive. This trade-off between the accuracy of the surrogate model and the cost of its construction presents a difficult decision for a design team, because additional sampling of the high-fidelity model does not guarantee any appreciable increase in surrogate accuracy.

This thesis presents a new method of surrogate synthesis that uses sampled data from a high-fidelity model paired with a low-fidelity model with space mapping techniques. This surrogate construction method is derived and explained in detail, and several applications of this methodology are presented. Chapter IV presents three conceptual applications of this theory using fabricated analytical models that were useful in the development and debugging of the space mapping algorithm presented in Chapter III. In Chapter V, this surrogate formation method is employed using tools made available by the Multidisciplinary Science and Technology Center (MSTC) within AFRL/RQ, and the results are compared with established surrogate construction techniques to determine what benefits, if any, this new method might offer a design team.

On a practical note, much of the literature pertaining to the subjects of surrogate construction and space mapping refer to the analytical model(s) employed in general terms because the methods themselves are meant to be applicable in a very general sense. As a result, the terminology and symbolism used in the derivation of these methods can be confusing. A section in the Appendix of this document serves to clarify the nomenclature. This section has been partitioned so the variables and their definitions match the appropriate sections of Chapters II and III to provide a more useful resource to the reader when navigating this document.

II. Background and Theory

To set the stage for the development of this new method for surrogate construction, the commonly accepted practices of surrogate modeling are reviewed. The following sections provide quick summaries of techniques that are applied in one form or another in the methodologies of this research. This chapter ends with a summary of a space mapping technique available in the literature which forms the foundation for much of this research.

Least-Squares Projections

In 1795, Friedrich Gauss developed a process for determining the best fit of linear coefficients to a general set of data, and this process has profoundly impacted most every scientific field since [2]. This process is known as the least-squares fitting of data, and is accomplished through the use of a projection matrix. Least-squares is still a popular method for approximating solutions to over-determined systems. An over-determined system of equations occurs when the number of equations is greater than the number of unknowns. The process of least-squares is often used to fit a linear model to data in the presence of noisy measurements, where the random noise on the measurements prohibits the solution to the system of equations through simple linear algebra.

The least-squares process is important to this research because it is a simple and powerful tool for fitting data to a prescribed analytical form. This is crucial to the space mapping process in Chapter III. To better understand the least-squares process, consider the following set of data developed as an illustrative example in Table 1.

Table 1. Noisy Measurements

$\bar{x}$	$\bar{y}$	$\bar{z}$
0.1517	0.3244	0.3772
0.1079	1.5886	0.8641
1.0616	0.6224	1.142
1.5583	1.0571	3.4421
1.868	0.3313	1.5041
0.2598	1.204	2.1154
1.1376	0.5259	1.7403
0.9388	1.3082	3.2392
0.0238	1.3784	0.5586
0.6742	1.4963	2.0559

The z data were generated according to the following equation, which is a simple linear combination of the x and y variable values, plus a random element to simulate the presence of noise in the measurement.

$$z = \bar{C}(1)x + \bar{C}(2)y + \{\text{random} \in [-1, 1]\} \quad \text{where} \quad \bar{C}(1) = \bar{C}(2) = 1 \quad (1)$$

If there were no noise in any of these measurements, then the coefficients $\bar{C}(1)$ and $\bar{C}(2)$ in the equation above could be deduced using Gaussian elimination. In the presence of noise however, not all of the equations are in perfect alignment with the assumption that z is a linear combination of x and y . To approximate the linear coefficients in the presence of noise, the least-squares process can be used.

The first step is to cast the problem in matrix form, as shown in Equation 2. Since there are coefficient values for this system of equations that will exactly fit the data collected, a least-squares fit determines the coefficient values that minimize the error between the gathered data and the resulting linear model. This can be seen in Figure 3, where the data points are shown as blue circles, and the linear approximation is shown as the mesh surface. The lines connecting the data points to the linear approximation represent the error between the data and the linear model:

$$\mathbf{A} \bar{\mathbf{C}} = \bar{\mathbf{z}} \quad \text{where} \quad \mathbf{A} = [\bar{\mathbf{x}} \quad \bar{\mathbf{y}}]. \quad (2)$$

The least-squares fitting of data works through the use of a projection matrix. The derivation of this projection matrix is an interesting topic for study because in its derivation this matrix is shown to minimize the total error between the gathered data and the linear model through a single analytical process (rather than requiring some number of iterations). See [3] for a detailed derivation of the projection matrix and greater detail into the least-squares process. The projection matrix is formed according to the equation below.

$$\mathbf{P}_{\mathbf{A}} = (\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T \quad (3)$$

Once the projection matrix is formed, the linear coefficients are determined through the following expression:

$$\bar{\mathbf{C}} = \mathbf{P}_{\mathbf{A}} \bar{\mathbf{z}}. \quad (4)$$

The linear model derived by this least-squares approach for the system shown in Table 1 is

$$z = \bar{C}(1)x + \bar{C}(2)y = 1.084x + 0.933y. \quad (5)$$

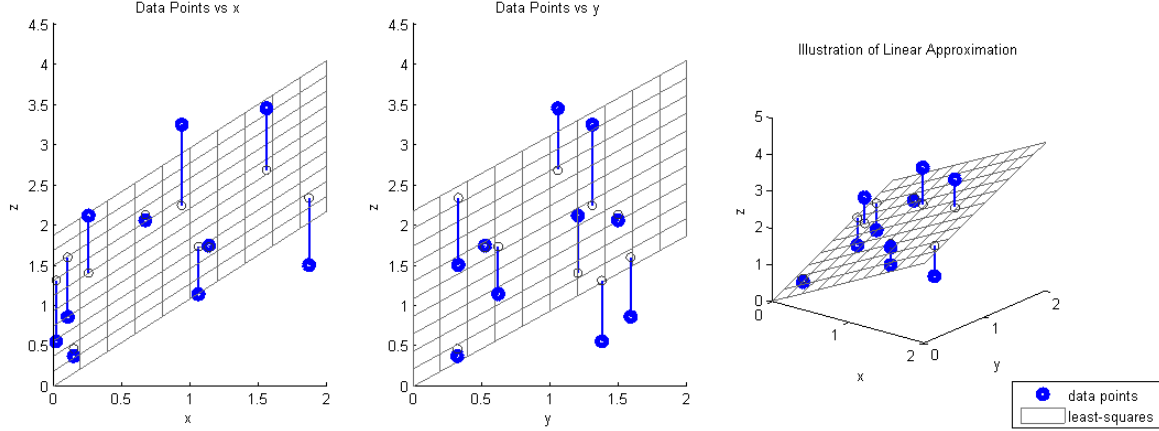


Figure 3. Illustration of Least-Squares Approximation

The least-squares process is not limited in the number of variables (or columns in the **A** matrix) which allows for least-squares fitting of data that are functions of many variables. Many of the mathematical forms employed in the space mapping implementation make use of a least-squares fit to determine the parameters within a particular **P**. The least-squares approach is beneficial to this research because the minimization of error between the data and the approximate model is conducted in a single step. Methods for non-linear least-squares approximations exist, but these methods require some number of iterations to determine the coefficients and powers which minimize the error between the data and the approximation.

Polynomial Response Methodology

In cases where access to data is limited, surrogate models can be generated from the existing data to make estimates at unknown locations in the design space by fitting a polynomial response through a least-squares projection. This type of surrogate technique is a subset of Polynomial Response Methodology (PRM). In this process, the surrogate model for the true response is represented by a linear combination of certain quantities. These quantities are derived from the inputs to the model being approximated, and the form for each quantity is left to the user to decide. If the coefficients in the polynomial form are determined in a least-squares sense, then the

polynomial must be linear with respect to these quantities. The choice of the form for each quantity allows the polynomial to be nonlinear with respect to the design variables. From this point forward, any reference to the order of a polynomial form is always with respect to the design variables [4].

In general, a polynomial of the k^{th} degree can be fit to a collection of data, provided there is enough data to produce an over-determined system of equations. The appropriate degree of polynomial to implement is often decided through experience or can be determined through the use of estimated error variances. The mathematical form of this polynomial can be represented as

$$R(\bar{x}) = \sum_{i=1}^p \left[\sum_{j=1}^k (C_j(i) \bar{x}(i)^j) \right] + C_0 \quad (6)$$

where $R(\bar{x})$ is the response of the polynomial, p is the number of variables in the system, and k is the polynomial degree being applied. Using the same variables from Table 1 and assuming a first-order polynomial fit results in the equation:

$$C_1(1)x + C_1(2)y + C_0 \quad (7)$$

while a second-order fit results in the equation:

$$C_2(1)x^2 + C_2(2)y^2 + C_1(1)x + C_1(2)y + C_0. \quad (8)$$

The advantage of the least squares process is the ability to determine the best coefficient values for a k^{th} order polynomial in a single linear algebra operation (see Equation 4). For polynomial degrees of two and higher, the $\mathbf{A}$ matrix shown in Equation 2 needs to be expanded to include the elements of the data raised to the appropriate power. A more general form for this matrix is

$$\mathbf{A} = \left[\begin{array}{cc|cc|ccc} \bar{x}_1(1)^k & \bar{x}_p(1)^k & \bar{x}_1(1)^{k-1} & \bar{x}_p(1)^{k-1} & \dots & \bar{x}_1(1) & \bar{x}_p(1) & 1 \\ \downarrow & \dots & \downarrow & \dots & \downarrow & \downarrow & \dots & \downarrow \\ \bar{x}_1(q)^k & \bar{x}_p(q)^k & \bar{x}_1(q)^{k-1} & \bar{x}_p(q)^{k-1} & \dots & \bar{x}_1(1) & \bar{x}_p(1) & 1 \end{array} \right] \quad (9)$$

and the dimensions of the matrix are $[q \times (k \cdot p + 1)]$, where q is the number of observed data points and p is the number of variables in the function being approximated. For a second-order ($k = 2$) fit to the data in Table 1, the $\mathbf{A}$ matrix is shown in Equation 10. The projection matrix for

any given polynomial degree can be calculated by forming the appropriate $\mathbf{A}$ matrix using Equation 9 to insert into Equation 3.

$$\mathbf{A} = \begin{bmatrix} \left| \begin{array}{cc} \bar{x}(1)^2 & \bar{y}(1)^2 \\ \downarrow & \downarrow \\ \bar{x}(10)^2 & \bar{y}(10)^2 \end{array} \right| \left| \begin{array}{cc} \bar{x}(1) & \bar{y}(1) \\ \downarrow & \downarrow \\ \bar{x}(10) & \bar{y}(10) \end{array} \right| \left| \begin{array}{c} 1 \\ \downarrow \\ 1 \end{array} \right| \end{bmatrix} \quad (10)$$

Users of polynomial response surfaces must be aware, unless the observed data were generated from a polynomial function, there is model error associated with any polynomial fit. This is due to the fact the observed data will almost never align perfectly with a polynomial formula, either because the process under consideration is more complex than the polynomial form will allow, or because of the presence of noise in the data. The only knowledge available for the process being modeled is the observed data points, and the error between these data points and the polynomial response. This error is called the actual absolute error [5], and is determined by

$$e_{act} = |\bar{z} - \hat{z}| \quad (11)$$

where $\bar{z}$ is a vector of observed data points and $\hat{z}$ is a vector of polynomial responses for the same variable values.

The true error at a new position cannot be known without first observing another data point at the same position. Making this new observation is impractical, since the requirement for a polynomial response surface infers an unwillingness or inability to observe additional data on the part of the design team. Estimation of the error at a new position is possible, however, by calculating the standard error (or the square root of prediction variance). This predicted error is given by the equation [5]:

$$e_{es}(\bar{x}) = \pm \sigma \sqrt{f^T(\bar{x}) (\mathbf{A}^T \mathbf{A})^{-1} f(\bar{x})} \quad (12)$$

where σ is the standard deviation of the error between observed data and polynomial approximations and $f(\bar{x})$ is a vector of variable values raised to the appropriate power.

$$f^T(\bar{x}) = [\bar{x}(1)^k \cdots \bar{x}(p)^k, \bar{x}(1)^{k-1} \cdots \bar{x}(p)^{k-1}, \rightarrow \bar{x}(1)^2 \cdots \bar{x}(p)^2, \bar{x}(1) \cdots \bar{x}(p), 1]. \quad (13)$$

σ is calculated using the Equation 14, where $\bar{z}$ is the observed data, $\bar{R}$ is the polynomial response at the same positions as $\bar{z}$, q is the number of observed data points, and the quantity $(kp + 1)$

represents the number of coefficients in the polynomial approximation [5]. Note the standard error measure shown in Equation 12 is a localized value and therefore dependent upon the position being evaluated. This estimation of error at a point could be useful in choosing the best PRM surrogate model to apply at a given location.

$$\sigma = \sqrt{\frac{(\bar{z} - \bar{R})^T (\bar{z} - \bar{R})}{(q - (kp + 1))}} \quad (14)$$

Kriging

Kriging is an interpolation method developed in the field of geostatistics. Geoff Bohling, with the Kansas Geological Survey, describes kriging as “optimal interpolation based on regression against observed z values of surrounding data points, weighted according to spatial covariance values [6].” The process is heavily rooted in linear regression analysis, with the distinction of choosing the weighting associated with sampled data points in an optimal fashion. The general form for a linear regression estimator is

$$z^*(\bar{x}) - m(\bar{x}) = \sum_{\alpha=1}^q \lambda_{\alpha} [z(\bar{x}_{\alpha}) - m(\bar{x}_{\alpha})] \quad (15)$$

where $z^*(\bar{x})$ is the estimate of z at some location $\bar{x}$, $m(\bar{x})$ is the expected value at some location $\bar{x}$, q is the number of sample data points, λ_{α} is the weighting coefficient associated with each sample data point, $z(\bar{x}_{\alpha})$ is the z value at the sample location $\bar{x}_{\alpha}$, and $m(\bar{x}_{\alpha})$ is the expected value of z at the sample location $\bar{x}_{\alpha}$ [7]. Kriging uses this same form, but makes use of a covariance matrix constructed from the sample locations to choose weighting coefficients (λ_{α}) such that the variance of the estimate for $z^*(\bar{x})$ is minimized.

A covariance matrix is a collection of the covariances between the sample data points. Covariance is calculated between two random variables and is the degree to which one variable trends in the same direction as the opposing variable. A positive covariance value indicates an increase in one of the variables likely correlates to an increase in the other variable. The covariance matrix used in the kriging process is calculated for the same random variable, $z(\bar{x})$, to determine how the value of each sample point is correlated to the other sample points. The formula used to construct this covariance matrix is

$$K_{ij} = \text{cov}(z(\bar{x}_i), z(\bar{x}_j)) = \text{E}[(z(\bar{x}_i) - m(\bar{x}_i))(z(\bar{x}_j) - m(\bar{x}_j))] \quad (16)$$

where $\mathbf{K}$ is the covariance matrix, and E is the notation for the expected value of the expression enclosed in brackets [6].

The goal in choosing the weighting coefficients in the kriging process is to minimize the variance of the resulting estimate:

$$\sigma_z^2(\bar{x}) = \text{var} \{z^*(\bar{x}) - z(\bar{x})\}. \quad (17)$$

The expected value for the quantity $z^*(\bar{x}) - z(\bar{x})$ is zero, which is called the unbiasedness constraint [6]. The sample data points can be viewed as the sum of two components, a trend component ($m(\bar{x})$) and a residual component ($r(\bar{x})$):

$$z(\bar{x}) = r(\bar{x}) + m(\bar{x}). \quad (18)$$

The expected value for the residual component is also assumed to be zero. The covariance of the residual component is assumed to be stationary, which means the covariance between sample points is a function of lag, $\bar{h}$, but not position, $\bar{x}$.

$$\text{cov} \{r(\bar{x}), r(\bar{x} + \bar{h})\} = E \{r(\bar{x}) \cdot r(\bar{x} + \bar{h})\} = C_R(\bar{h}) \quad (19)$$

This residual covariance function is usually determined from the user's choice of correlation function and should ideally represent the residual component of the sample data points [6]. A correlation function is a model that represents the dependence of the sample data points $z(\bar{x})$ on the values in the location vector ($\bar{x}$), and the choice of this model is an important assumption, which affects the kriging estimates between sample data points [8].

There are many different variants of the kriging process, but an explanation of simple kriging suffices here. The main difference between the variants of kriging is the manner in which they treat the trend component of the sample data, $m(\bar{x})$ [7]. Simple kriging treats the trend component as a constant value equal to the average value of the sample data points.

$$m(\bar{x}) = \frac{\sum z(\bar{x})}{q} = m \quad (20)$$

Equation 15 can now be rewritten as

$$z^*(\bar{x}) = m + \sum_{\alpha=1}^q [\lambda_{\alpha} (z(\bar{x}_{\alpha}) - m)]. \quad (21)$$

The estimation error is then calculated as $z^*(\bar{x}) - z(\bar{x})$, and can be represented as a linear combination of the difference between the estimated residual and the sample residual at each sample location ($r^*(\bar{x}_{\alpha}) - r(\bar{x}_{\alpha})$ as α varies from 1 to q).

$$z^*(\bar{x}) - z(\bar{x}) = [z^*(\bar{x}) - m] - [z(\bar{x}) - m] \quad (22)$$

$$= r^*(\bar{x}) - r(\bar{x}) \quad (23)$$

$$= \sum_{\alpha=1}^q [\lambda_{\alpha} r(\bar{x}_{\alpha})] - r(\bar{x}) \quad (24)$$

Applying the mathematical rules governing the variance of a linear combination of random variables, the variance of the error can be calculated by:

$$\sigma_{error}^2 = \text{var} \{r^*(\bar{x})\} + \text{var} \{r(\bar{x})\} - 2 \text{cov} \{r^*(\bar{x}), r(\bar{x})\} \quad (25)$$

$$= \sum_{\alpha=1}^q \sum_{\beta=1}^q \lambda_{\alpha} \lambda_{\beta} C_R(\bar{x}_{\alpha} - \bar{x}_{\beta}) + C_R(0) - 2 \sum_{\alpha=1}^q \lambda_{\alpha} C_R(\bar{x}_{\alpha} - \bar{x}_{\beta}). \quad (26)$$

Equation 26 is the quantity minimized in the kriging process through the proper selection of values for λ . To minimize this quantity, the derivative of Equation 26 with respect to λ_{α} is taken and set to zero. The result is the following system of equations [6]:

$$\sum_{\beta=1}^q \lambda_{\beta} C_R(\bar{x}_{\alpha} - \bar{x}_{\beta}) = C_R(\bar{x}_{\alpha} - \bar{x}_{\beta}) \quad \text{for } \alpha = 1 \text{ to } q. \quad (27)$$

The values for the covariance between sample data points is determined using Equation 19, and the system of equations shown in Equation 27 can be written in matrix form as

$$\mathbf{K} \bar{\lambda} = \bar{k} \quad (28)$$

where $\mathbf{K}$ is the covariance matrix for the sample data, and $\bar{k}$ is the vector of covariance values between the sample data points and the estimation point. The weighting coefficients that minimize

the variance of the error estimate can then be found by solving for $\bar{\lambda}$:

$$\bar{\lambda} = \mathbf{K}^{-1} \bar{k}. \quad (29)$$

The estimate at a given location is given by Equation 21 with the appropriate weighting coefficients applied.

Kriging is a variant of an interpolation method, and as such the process shares the same properties common to all interpolation techniques. The accuracy of the estimate is dependent upon the locations of the sample points relative to the estimation point. Good estimates using any interpolation technique require a sufficient number of appropriately distributed sample points in the region of interest. If the sample point locations are biased or clustered, then an interpolation routine (kriging or otherwise) will not yield accurate estimates away from the sampled data [6].

Kriging does offer some distinct advantages over other interpolation techniques. The minimum variance of the estimation error at the sample data locations equals zero because the true value of $z(\bar{x})$ is known. As a result the kriging model will match the values of the sample locations perfectly, which is not true of other interpolation routines. Due to the manner in which the weights for each sample point are derived, kriging applies reduced weighting to points within data clusters as opposed to solitary data points. Since the kriging process involves minimizing the estimation error variance in Equation 26, a kriging model yields not only an estimate at a given location but also the variance of the estimation error at that location [6]. This allows for a kriging model to provide both an estimated value at a location in the design space, and a measure of how accurate the kriging model is at that location.

Space Mapping

The concept of space mapping was developed by Dr. John Bandler in 1993; space mapping has been used in various engineering applications over the years. Since its introduction, numerous teams and individuals have contributed to the advancement of the techniques employed in space mapping, and Bandler *et al.* have published a paper detailing many of these achievements [9]. As background for this thesis, this section will outline the space mapping technique listed in [10], which hereafter is referred to as the “traditional” space mapping approach. This space mapping approach has been used at Sandia National Laboratories in the implementation of a multifidelity analysis infrastructure [10].

The purpose of space mapping is to align the response of a low-fidelity (or coarse) model with the response of a high-fidelity (or fine) model over some region of the design space. This alignment is achieved by mathematically relating the inputs of the low-fidelity model, $\bar{x}_L$, to the inputs of the high-fidelity model, $\bar{x}_H$.

$$\bar{x}_L = \mathbf{P}(\bar{x}_H) \quad (30)$$

such that

$$R_L(\bar{x}_L) \approx R_H(\bar{x}_H). \quad (31)$$

R_L is the response of the low-fidelity model and is, through the variable relationship $\mathbf{P}$, a function of the high-fidelity design vector. R_H is the response of the high-fidelity model.

A traditional space mapping approach begins with an assumption of the form for the variable relationship, $\mathbf{P}$. This relationship can take on any formula deemed appropriate by the user. One of the simpler forms for $\mathbf{P}$ is for each low-fidelity design variable to be a linear combination of the high-fidelity design variables. This is the form assumed for this space mapping overview.

$$\begin{bmatrix} \bar{x}_L(1) \\ \vdots \\ \bar{x}_L(n) \end{bmatrix} = \begin{bmatrix} \phi_{1,1} & \cdots & \phi_{1,p} \\ \vdots & \ddots & \vdots \\ \phi_{n,1} & \cdots & \phi_{n,p} \end{bmatrix} \begin{bmatrix} \bar{x}_H(1) \\ \vdots \\ \bar{x}_H(p) \end{bmatrix} \longrightarrow \bar{x}_L = \Phi \bar{x}_H \quad (32)$$

One important aspect of the space mapping process employed at the Sandia National Laboratories is the ability to relate design vectors of different lengths. In Equation 32, the lengths of $\bar{x}_L$ and $\bar{x}_H$ are not required to be equal. The low-fidelity design vector has n elements, the high-fidelity design vector has p elements, and for the purposes of this research $p \geq n$.

The next step in the space mapping process is to query the high-fidelity model to obtain a response for various design variable values. Once these high-fidelity responses have been collected, the space mapping process calls for a minimization sequence. This step seeks the parameters within the mathematical relationship $\mathbf{P}$ that cause the response of the low-fidelity model to be approximately equal to the high-fidelity response. This minimization process is shown in the equation:

$$\min_{\mathbf{P}} J = \sum_{j=1}^q [R_H(\bar{x}_{H_j}) - R_L(\mathbf{P}(\bar{x}_{H_j}))]^2 \quad (33)$$

This minimization process seeks to minimize the total error between the queried high-fidelity responses and the responses from the low-fidelity model. The degrees of freedom in this

minimization process are the parameters within $\mathbf{P}$, while the design variables associated with each high-fidelity response are held constant. The output of this process is a collection of parameter values for $\mathbf{P}$ that most closely replicate the high-fidelity responses gathered using the low-fidelity model as a surrogate.

Once a reliable space mapping has been derived between the two models, it is possible to approximate the response of the high-fidelity model for a given $\bar{x}_H$ by executing the low-fidelity analysis using the variable relationship assumed in Equation 30. The low-fidelity model, in conjunction with the space mapping relationship, can then act as a surrogate analysis for the high-fidelity model. This can be advantageous if the expense of running the high-fidelity model precludes its use in an optimization scheme or design space exploration.

III. Methodology

Space mapping is the process of relating the variables in the low-fidelity model to the variables in the high-fidelity model such that the difference between the two models is minimized. The variable relationship $\mathbf{P}$ is a mathematical formula, and so a traditional space mapping approach makes an assumption as to what form this relationship takes. This research develops and explores an alternative method of space mapping which does not require an assumption of the particular form for $\mathbf{P}$. Rather, this modified space mapping algorithm allows for multiple variable relationships to be assumed through the fitting of mathematical models to data gathered in the new process. The resulting space-mapped surrogate models can be evaluated with respect to each model's ability to approximate the high-fidelity response. Thus, the best performing form for the variable relationship $\mathbf{P}$ can be determined without the need for multiple space mapping iterations.

Assumptions and Limitations

As in many engineering applications, the methods discussed in this research have been built upon certain assumptions. As a result, these methods are also limited in the scope of their application by the bounds of the assumptions made. A list of assumptions and their associated limitations have been compiled in Table 2. The first assumption, and likely the most important, is the one buried at the core of any space mapping process: these methods require a pair of analyses that calculate the same parameter at differing levels of fidelity and computational expense. Without this analytical pairing, the methods employed within this research will be impossible to implement.

An additional assumption to the space mapping process for this research is that all of the low-fidelity variable definitions are present in the high-fidelity variables. This assumes the high-fidelity model always has at least as much information about the system as the low-fidelity model. This assumption results in the concept of shared design variables, which are those design variable definitions (not the design variable values) that are common between the two models. The presence of shared design variables leads to the ordering of the high-fidelity design vector elements as follows:

$$\bar{\mathbf{x}}_H = \begin{bmatrix} \tilde{\mathbf{x}}_L \\ \tilde{\mathbf{x}}_H \end{bmatrix} \quad (34)$$

Table 2. List of assumptions and associated limitations

#	Assumption	Limitation
1	Appropriate pairing of similar high and low-fidelity models	Space mapping process is not applicable without this specific analytical pairing
2	All variables in the low-fidelity model are also present in the high-fidelity model	Exclusion of model-pairings where the low-fidelity model contains inputs not present in the high-fidelity model
3	The sampled points in the high-fidelity space capture the necessary information for a viable space mapping	Poor distribution of sample points will result in an inaccurate space mapping, which will yield inaccurate predictions of the high-fidelity response
4	Both the high and low-fidelity design spaces are continuous	Exclusion of models with discontinuous design spaces

where $\tilde{x}_L$ is the vector of shared design variables and $\tilde{x}_H$ is the vector of additional design variables present in the high-fidelity model (but not present in the low-fidelity model). Note that the low-fidelity design vector $\bar{x}_L$ is assumed to have the same dimensions as the shared design variables $\tilde{x}_L$ as defined in Equation 34. From a practical standpoint, the presence of the low-fidelity variables in the high-fidelity design vector greatly simplifies a step in the optimization process. This optimization process is discussed in further detail in the section devoted to the modified space mapping algorithm.

A key component in any space mapping approach is the sampling of the high-fidelity design space. Only a limited number of data points are taken from this design space due to the relative expense of running the high-fidelity model. Each data point is crucial to the space mapping process because it represents information about the high-fidelity design space that is passed to the low-fidelity design space through the space mapping. The goal in the sampling of the high-fidelity design space is to gain the necessary information for the least number of points. The validity of any given space mapping is, in a sense, a reflection of how well the sampling process gathered the necessary information from the high-fidelity design space. As a result, there is an ever-present balancing act between the number of points that are affordable to the design team and the number of points required for a sufficiently accurate space mapping. Without prior knowledge of the design space, the goal of attaining information as cheaply as possible is best served through an even distribution of sampled points within set variable boundaries. Two common methodologies for achieving an even distribution of samples are orthogonal arrays and latin hypercube sampling. Orthogonal arrays have the advantage that they ensure coverage of the sample space boundaries,

while latin hypercube spacing adds a degree of randomness to the sampling that helps prevent aliasing effects. This research uses the latter method, which is implemented through an algorithm included in the **dace** MATLAB extension [8]. An excellent description of the latin-hypercube sampling methodology (complete with MATLAB code) can be found in [1].

The gradient-based optimization process required in this space mapping approach carries an assumption that the design space of each model is continuous in the region of interest. In general, a gradient-based optimization scheme uses first or second-order derivative information to determine the direction of the optimum point. These mathematical foundations make gradient-based optimizers ill-equipped to handle instantaneous changes in first or second-order derivatives. Certain gradient-based optimization routines may include algorithms to help optimize a solution in the presence of discontinuities, but no such tools were implemented in this research. As a result, the methods to be outlined in this thesis are employed using continuous analytical models. A genetic algorithm in place of a gradient-based optimizer might eliminate this limitation, but this alternative was not explored in this research.

Modified Space Mapping Approach

The traditional space mapping approach assumes a form for the variable relationship and uses the optimization sequence outlined in Equation 33 to determine the parameters within $\mathbf{P}$. Given a sufficiently pliable relationship for $\mathbf{P}$, this process has yielded viable space mappings [10]. The alternative space mapping approach developed in this research changes the algorithm such that the variable relationship $\mathbf{P}$ is not chosen until after the minimization process has taken place. Not assuming a form for $\mathbf{P}$ going into the minimization sequence requires modifications to the space mapping algorithm and the objective function for the minimization. Rather than perform a single minimization to reduce the total error between the sampled high-fidelity data and the space-mapped low-fidelity model, this new process performs a minimization sequence for each high-fidelity data point taken.

While the traditional approach allows the optimizer to change the parameters within an assumed form of $\mathbf{P}$, the modified approach allows the optimizer to directly change the low-fidelity variables themselves. Since the optimizer changes the values of the low-fidelity variables, the output of each minimization sequence is a vector of low-fidelity variables associated with the vector of high-fidelity variables that resulted in a particular data point. Each data point has a minimization sequence to produce an optimal low-fidelity design vector (labeled $\bar{x}_L^*$) that

minimizes the difference between fidelity levels. The accumulation of these low-fidelity design vectors and their associated high-fidelity design vectors is the data used in the space mapping process. The objective function for each of these minimization sequences is

$$\min_{\bar{x}_L} J = [R_H(\bar{x}_H) - R_L(\bar{x}_L)]^2. \quad (35)$$

In addition to the optimal low-fidelity design vector, the optimization process shown in Equation 35 also yields the residual difference between the high and low-fidelity responses. There will be instances in which the low-fidelity model is incapable of attaining the magnitude of the response seen in the high-fidelity model over the region of the design space in which $\mathbf{P}$ is derived. In these circumstances, the vector of residual differences will be nonzero and each vector element can be associated with the $\bar{x}_H$ that resulted in the appropriate R_H . The residual differences can then be approximated using the same techniques used to fit forms to the low-fidelity variables. The residual differences are calculated at the end of each minimization sequence by the equation

$$\Delta R = R_H(\bar{x}_H) - R_L(\bar{x}_L^*). \quad (36)$$

This information is used to construct a model of the local difference between fidelity levels as a function of the high-fidelity design variables. In keeping with the nomenclature for space mapping techniques, the model for the local difference between fidelity levels is referred to as $\Delta\mathbf{P}$. The output of this local difference model is then superimposed upon the space-mapped low-fidelity response to better approximate the high-fidelity response, as shown in Figure 4.

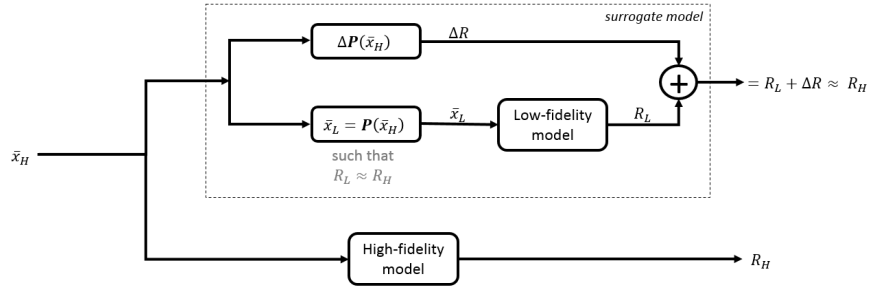


Figure 4. Illustration of a space mapped surrogate model receiving the same inputs as the high-fidelity analysis to produce an approximation of the high-fidelity response

The collection of low and high-fidelity design vectors shown in Table 3 can be interpreted as a collection of outputs (the elements of the low-fidelity design vector and the residual difference) for a given collection of inputs (the high-fidelity design vectors). If each variable in the collection of optimal low-fidelity design vectors is considered independently, then an analytical form (either a polynomial or kriging surface) can be fit for that specific low-fidelity design variable. Since the minimization sequences changed the low-fidelity design variables rather than the parameters of an assumed $\mathbf{P}$, the form of the space mapping relationship for each low-fidelity design variable can be either chosen by the user or determined by a comparisons of total error values for a number of relationships. In this modified version of space mapping, the form of the variable relationship is a post-processing task that does not require additional executions of either the high or the low-fidelity analysis for new forms of $\mathbf{P}$ to be fit to the data.

Table 3. Collection of data from the modified space mapping algorithm

#	outputs			inputs		
	low-fidelity design vectors			high-fidelity design vectors		
1	$\bar{x}_{L_1}^*(1)$	$\cdots$	$\bar{x}_{L_1}^*(n)$	$\Delta\bar{R}(1)$	$\bar{x}_{H_1}(1)$	$\cdots$ $\bar{x}_{H_1}(p)$
2	$\bar{x}_{L_2}^*(1)$	$\cdots$	$\bar{x}_{L_2}^*(n)$	$\Delta\bar{R}(2)$	$\bar{x}_{H_2}(1)$	$\cdots$ $\bar{x}_{H_2}(p)$
$\downarrow$		$\downarrow$		$\downarrow$		$\downarrow$
q	$\bar{x}_{L_q}^*(1)$	$\cdots$	$\bar{x}_{L_q}^*(n)$	$\Delta\bar{R}(q)$	$\bar{x}_{H_q}(1)$	$\cdots$ $\bar{x}_{H_q}(p)$

Modified Space Mapping Algorithm

The process for the modified space mapping approach is divided into three main tasks, with each task consisting of several steps: sample the high-fidelity model, execute the minimization sequences to gather the data, and form (and evaluate) some number of variable relationships. Each task is composed of numerous steps which are outlined in this section. The end result of this space mapping process is a variable relationship, $\bar{x}_L = \mathbf{P}(\bar{x}_H)$, and an approximation for the residual difference, $\Delta\mathbf{P}(\bar{x}_H)$, that allows the low-fidelity analysis to approximate the high-fidelity response (within the bounds of the design space over which the space mapping was derived). A flowchart showing the three tasks within the space mapping algorithm is shown in Figure 5.

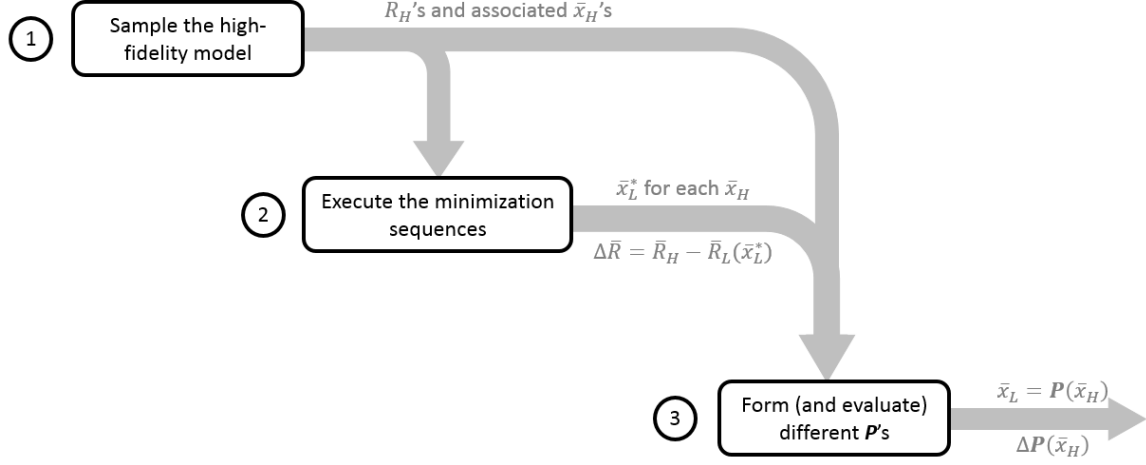


Figure 5. Modified space mapping algorithm, shown with inputs and outputs for each task

Task 1: Sampling the High-fidelity Design Space

Any space mapping process begins with a sampling of responses from the high-fidelity design space. This sampling process is the means by which information about the higher fidelity response as a function of the design variables is gathered. The distribution of these sampled points within a subset of the overall design space is vitally important to the accuracy of the variable relationship produced. The goal for the sampling process is to capture the significant design space contours of the high-fidelity response for the least number of points. This goal is best served by evenly distributing the sampled points within a subspace of the overall design space.

The first step in the sampling process (step 1.1 in Figure 6) is to define the bounds of this subspace. This is accomplished by setting upper and lower boundaries for the high-fidelity design variables. This is followed by sampling points within this defined subspace (step 1.2), where the sampling function will assign variable values within the bounds set by the user. The importance of the sampling process cannot be overstated due to the desire to be as efficient as possible in the gathering of information from the high-fidelity subspace. In the hopes of efficiently sampling the subspace, the algorithm implemented in this research makes use of the `lhsu()` function that is included in the `dace` MATLAB software extension [8]. This function generates a group of sample high-fidelity design vectors using a latin hypercube sampling methodology (the `u` character in `lhsu` stands for “uniform”, which refers to its use of the uniform distribution).

Once some number of evenly-distributed design vectors have been generated (the number of samples is designated q), step 3.1 is to generate a high-fidelity response for each sampled design

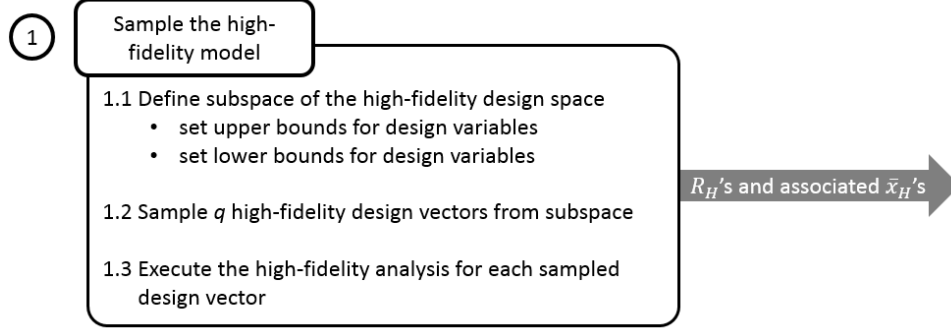


Figure 6. Steps that compose the first task in the modified space mapping process

vector from the previous step. The number of samples required to obtain a reliable space mapping will change based on the particular models being used and is likely dependent upon numerous factors, some of which are the variability of the high-fidelity design space in the region of interest and the *relative* variability between the low and high-fidelity model within the shared design space. The number of sampled points required for a reliable space mapping is a reflection of how much information from the higher fidelity is absent at the lower fidelity level. The exact characterization of this behavior was not investigated in this research. The number of sampled points required in the application of this process is noted, but there was no research towards predicting the appropriate number of samples. If this theory proves to be useful in a design environment, then the prediction of required sample points would be a relevant area for future work.

Task 2: Execute the Minimization Sequences

Once q number of points have been sampled from the high-fidelity design space, the minimization sequences can begin. The first step within this task is the scaling of the high-fidelity design variables through a process known as standardization (step 2.1). The purpose of scaling the design variables prior to the minimization sequence is to aid the optimizer in the search for $\bar{x}_L^*$. The scaling of the design variables alleviates some of the problems that gradient-based optimizers face when dealing with variables of significantly different magnitudes. A discussion on the merits of scaling the design variables, complete with a mathematical explanation for the benefits of scaling in an optimization context, can be found in [11]. An illustration of the steps required in Task 2 is shown in Figure 7.

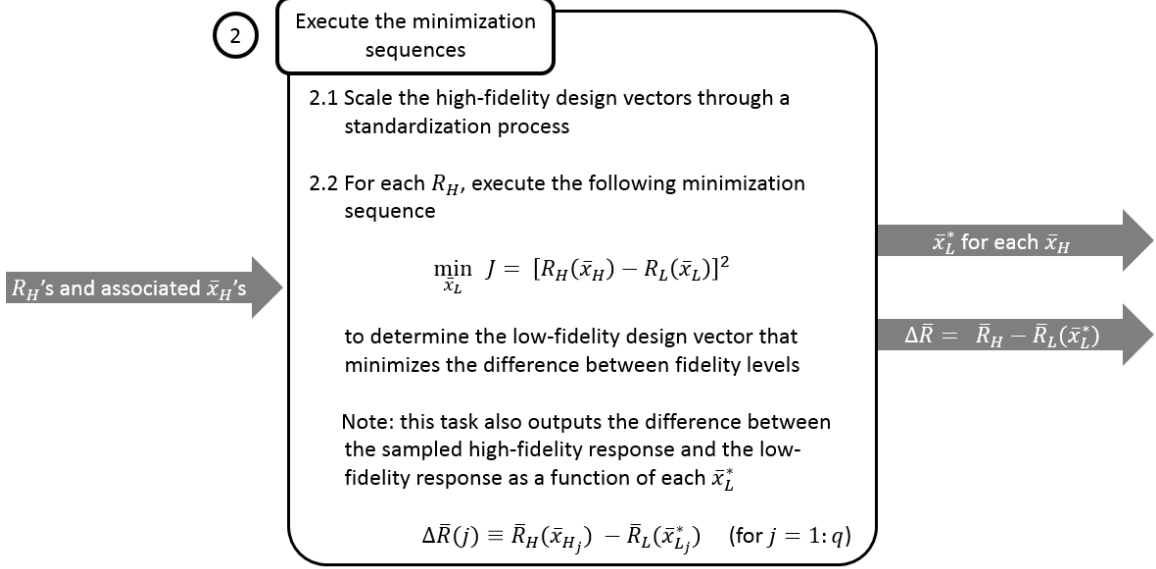


Figure 7. Steps that compose the second task in the modified space mapping process

The scaling for this space mapping approach is accomplished using a process known as standardization (or normalization). This process is widely used in the field of regression analysis and allows for the different orders of magnitude in the variables to be adjusted to a “standardized” scale. If the variables are scaled to a standardized order of magnitude, the optimizer can more easily determine the correct direction in the design space for the minimization of the objective function. The formula for the standardization of a given x value pulled from a sample population of x values is shown below

$$x_{scaled} = \frac{x - \mu_x}{\sigma_x} \quad (37)$$

where μ_x is the expected value of x and σ_x is the standard deviation of the sample. The expected value of x is calculated as the average value of the sample population.

The minimization sequences within Task 2 only require the scaling of the shared design variables. This is due to the fact that the objective function for each minimization (see Equation 35) only requires the low-fidelity design variables. The additional design variables at the higher fidelity level were required to calculate R_H , but these variables are not explicitly present in the objective function. In Task 3A, where the variable relationship $\mathbf{P}$ is assumed to be a k^{th} degree polynomial, the algorithm does require all of the high-fidelity design variables to be scaled, so the process diagram shown in Figure 7 includes the scaling of all the design variables within Task 2.

The purpose of the minimization sequences in step 2.2 is to gather the data needed to fit $\mathbf{P}$ and $\Delta\mathbf{P}$ in subsequent steps. A distinct minimization process will take place for each high-fidelity data point sampled, which means there will be q minimizations in total. The computational cost associated with this space mapping process up to this point consists of the cost of executing q high-fidelity analyses (Task 1) plus however many low-fidelity analysis calls are required in the minimization sequences (Task 2). The first two tasks in the modified space mapping algorithm represent the bulk of the computational expense. At the completion of the second task, each sampled high-fidelity design vector has an associated low-fidelity design vector and residual difference between fidelity levels. The pairings of each low-fidelity design vector and residual difference with the appropriate high-fidelity design vector constitute the dataset to be used in the final task.

Each minimization sequence requires a starting position in the shared design space (a starting location within the design space is required of any gradient-based optimization scheme). Earlier an assumption was made that the two analyses, while at different levels of fidelity, share some degree of similarity in the shared design space. Following this assumption, the starting point in the low-fidelity design space is set to the same location of the shared design space where the high-fidelity data point was sampled. The degrees of freedom for each minimization process are the *scaled* low-fidelity design variables, so the starting location for each minimization sequence consists of the scaled shared design variables in $\bar{x}_H$. The use of scaled variables in the optimization process requires the code that calculates the objective function to reverse the scaling process before executing the low-fidelity analysis.

At the completion of Tasks 1 and 2, the user has a collection of high-fidelity design vectors paired with low-fidelity design vectors and residual differences. These pairings are the data used in the final task of the space mapping algorithm. At this point, the user may choose the form of the variable relationship $\mathbf{P}$ that is believed to best fit the data collected. Three options are explored in this research, as shown in Figure 8.

Task 3A defines $\mathbf{P}$ as a k^{th} degree polynomial whose coefficients can be determined in a least-squares sense. Task 3B assumes a nonlinear polynomial form for $\mathbf{P}$ that requires an optimization sequence to determine the correct parameter values within $\mathbf{P}$. Task 3C assumes a highly nonlinear form for $\mathbf{P}$ where each low-fidelity design variable is represented by a kriging approximation of the variable data gathered in Tasks 1 and 2. With the exception of the calculation of total error, the fitting of the data to a specific variable relationship form does not

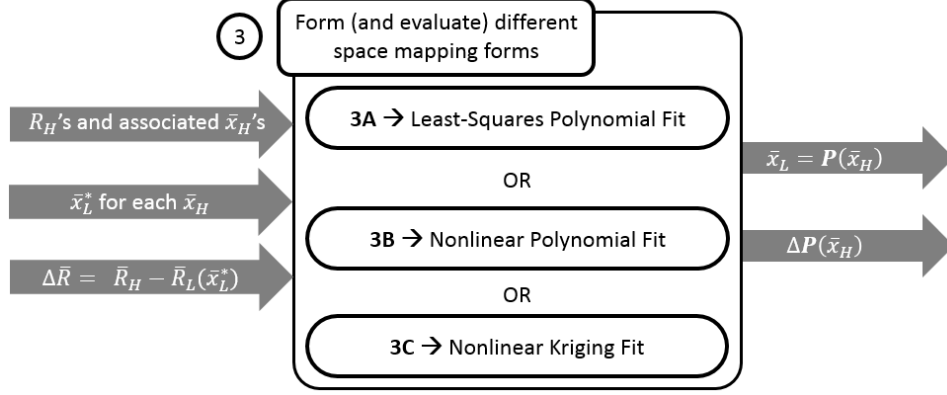


Figure 8. Steps that compose the final task in the modified space mapping process

require any additional executions of either the high or low-fidelity models. The total error calculation requires q additional low-fidelity analysis calls, as discussed in this section.

Task 3A: Least-Squares Polynomial Fit

The first step in the fitting of a least-squares polynomial to the data is to construct the projection matrix, $\mathbf{P}_A$. This requires the formation of the appropriate A matrix for the user-selected polynomial degree. The projection matrix is calculated using Equation 3 and is used in step 3.2 to determine the coefficients for the polynomial form. The vector $\bar{C}$ is comprised of coefficients for each variable raised to each power as well as an offset, which means this vector has a length of $(kp + 1)$.

$$\bar{C}_i = \mathbf{P}_A \begin{bmatrix} \bar{x}_{L_1}^*(i) \\ \vdots \\ \bar{x}_{L_q}^*(i) \end{bmatrix} \quad (\text{for } i = 1 : n) \quad (38)$$

This vector is also specific to the low-fidelity design variable currently being related to the collection of high-fidelity design vectors, which means the complete space mapping includes n number of $\bar{C}$ vectors. These steps are summarized in the Figure 9.

The second output from Task 2 is the vector $\Delta\bar{R}$ containing the residual differences between the sampled high-fidelity responses and the low-fidelity responses for the design variables determined in step 2.2. These residual differences are also fit to a polynomial approximation equivalent to the ones formed for each low-fidelity design variable. This results in a vector of

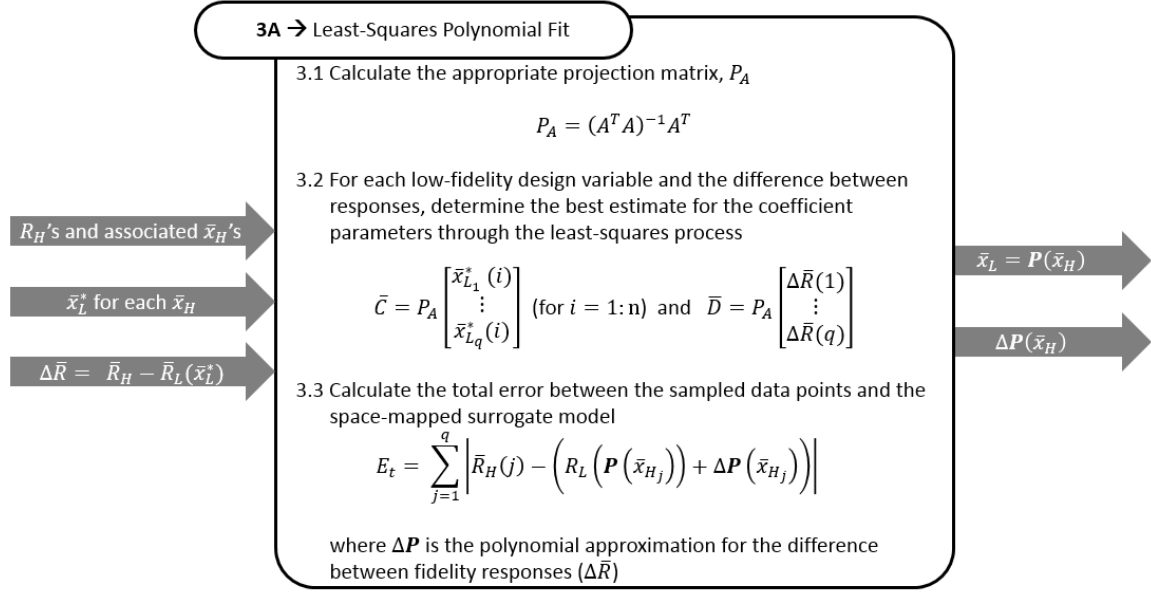


Figure 9. Steps specific to the least-squares polynomial form for $\mathbf{P}$

coefficients (defined as $\bar{D}$) found in the same manner as Equation 38.

$$\bar{D} = \mathbf{P}_A \begin{bmatrix} \Delta \bar{R}(1) \\ \vdots \\ \Delta \bar{R}(q) \end{bmatrix} \quad (\text{for } i = 1 : n) \quad (39)$$

The final step in Task 3A is the calculation of total error between the sampled high-fidelity data points and the response of the space mapped surrogate model. The error calculation is optional, but recommended for cases where multiple mapping forms are applied to the same dataset of associated low and high-fidelity design variables (from Tasks 1 and 2). The total error value is a measure of how well the resulting surrogate model approximates the known information from the high-fidelity design space and can be used to determine the best performing surrogate model. It is important to note the surrogate model with the lowest total error is *assumed* to be the best performer, but the user will never definitively know which surrogate performs best without an extensive interrogation of the high-fidelity design space. The formula to calculate the total error is shown below

$$E_t = \sum_{j=1}^q \left| \bar{R}_H(j) - \left(R_L \left(\mathbf{P}(\bar{x}_{H_j}) \right) + \Delta \mathbf{P}(\bar{x}_{H_j}) \right) \right| \quad (40)$$

where $\Delta \mathbf{P}$ is the polynomial approximation of the residual differences ($\Delta \bar{R}$) as a function of the high-fidelity design variables and the polynomial coefficients $\bar{D}$.

Task 3B: Nonlinear Polynomial Fit

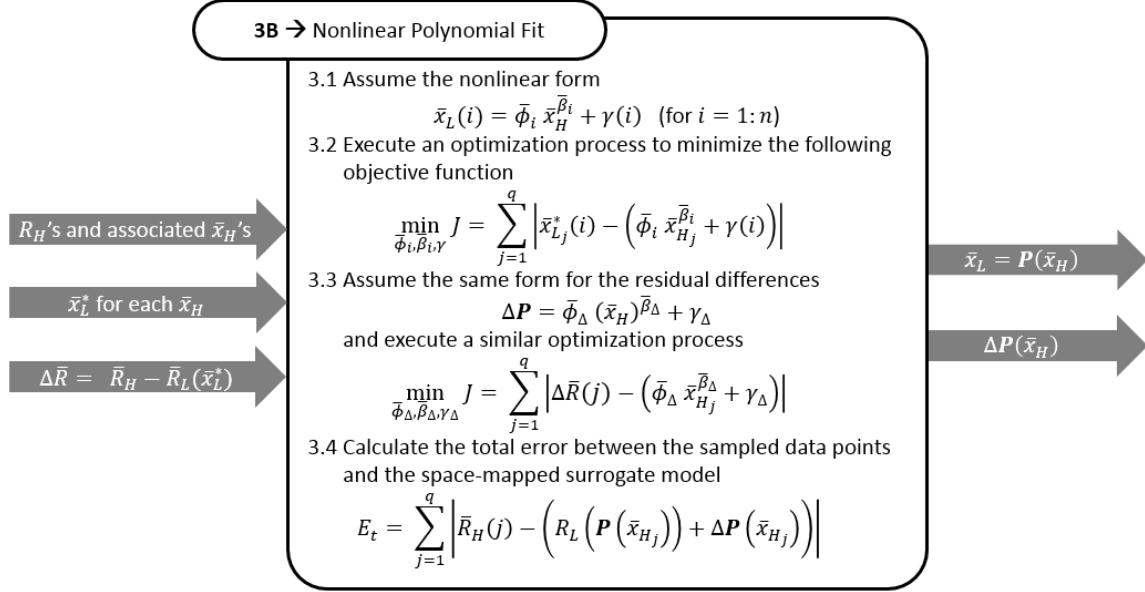


Figure 10. Steps specific to the nonlinear polynomial form for $\mathbf{P}$

While the steps in Figure 10 depict a particular nonlinear form, the process shown to determine the parameters of $\mathbf{P}$ is not specific to any given mathematical form. The form chosen for this research is a mapping formula recommended in [10],

$$\bar{x}_L(i) = \bar{\phi}_i \bar{x}_H^{\bar{\beta}_i} + \gamma(i) \quad (41)$$

where i is the element of the low-fidelity design vector, $\bar{\phi}_i$ is a vector of coefficients, $\bar{\beta}_i$ is a vector of powers, and $\gamma(i)$ is a scalar offset. Once a form for $\mathbf{P}$ has been determined, the next step is to determine the parameters of $\mathbf{P}$ using an optimization process. An optimizer seeks to find the parameter values of $\mathbf{P}$ that minimize the following objective function:

$$\min_{\bar{\phi}_i, \bar{\beta}_i, \gamma} J = \sum_{j=1}^q \left| \bar{x}_{Lj}^* - \left(\bar{\phi}_i \bar{x}_{Hj}^{\bar{\beta}_i} + \gamma(i) \right) \right| \quad (42)$$

where $\bar{x}_{L_j}^*(i)$ is the low-fidelity design variable found in Task 2 that is associated with the $\bar{x}_{H_j}$ sampled in Task 1.

The same relational form is then applied to the residual difference vector to derive a formula for $\Delta \mathbf{P}$. The difference between fidelity responses is approximated by the formula:

$$\Delta \mathbf{P} = \bar{\phi}_\Delta \bar{x}_H^{\bar{\beta}_\Delta} + \gamma_\Delta \quad (43)$$

The parameters of this relationship are determined through a similar optimization process as was shown previously.

$$\min_{\bar{\phi}_\Delta, \bar{\beta}_\Delta, \gamma_\Delta} J = \sum_{j=1}^q \left| \Delta \bar{R}(j) - \left(\bar{\phi}_\Delta \bar{x}_{H_j}^{\bar{\beta}_\Delta} + \gamma_\Delta \right) \right| \quad (44)$$

Once the parameters for $\mathbf{P}$ for each low-fidelity design variable and the parameters for the residual difference approximation $\Delta \mathbf{P}$ have been determined, the same total error calculation shown in Task 3A is performed (Equation 40).

Task 3C: Nonlinear Kriging Fit

The final space mapping form investigated in this research constructs a kriging model for each low-fidelity design variable as well as for the residual error between fidelity levels. In order to define a kriging model for a low-fidelity design variable, a vector is formed of the appropriate element in $\bar{x}_L^*$ from all of the design vectors output from Task 2. Going back to the symbolic data in Table 3, this vector is a single column of outputs. These vector elements are treated as the output from some model, and each element is some unknown function of the high-fidelity design variables that are the inputs in Table 3. This unknown function that relates the vector elements to the high-fidelity design vectors is modeled with a kriging approximator, resulting in a kriging model that can predict the value of that specific element in $\bar{x}_L^*$ as a function of $\bar{x}_H$. Each low-fidelity design variable (or each column of outputs in Table 3) has a kriging model associated with it, as will the residual difference between fidelity levels (as shown in Figure 11).

Due to the nature of kriging approximations, the step for calculation of the total error is removed from the implementation of this particular space mapping form. Kriging approximations have the unique property of being exact at the sampled points. This means that at each of the sampled high-fidelity design vectors, the resulting space-mapped low-fidelity design vector is exactly equal to $\bar{x}_{L_j}^*$ and the approximation for the residual difference between fidelity levels is

exactly equal to $\Delta \bar{R}(j)$. Thus, the total error calculation shown in Equation 40 is always zero for a kriging model implementation.

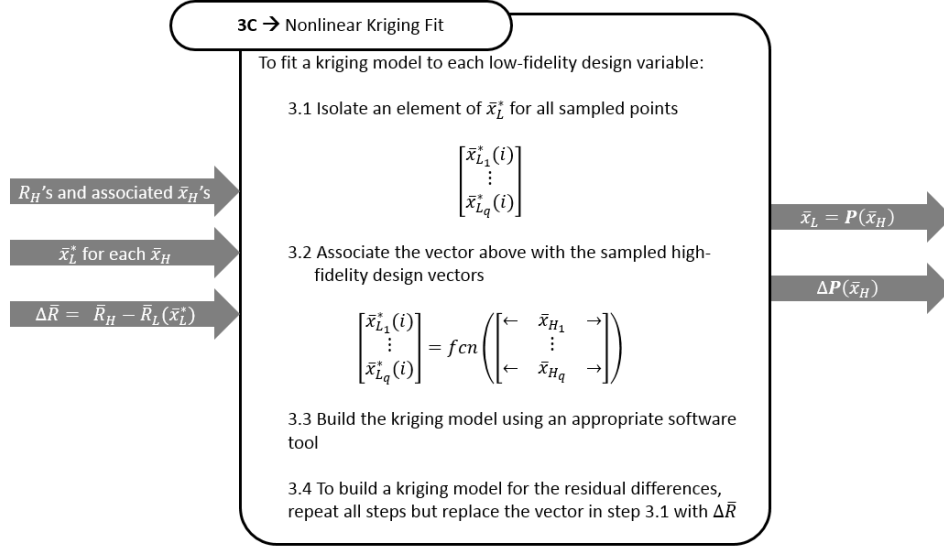


Figure 11. Steps specific to the nonlinear kriging space mapping form

While this makes a comparison between a kriging space mapped surrogate to other surrogates impossible using a total error calculation, the lack of a comparison is likely inconsequential. The kriging process requires a relatively large number of sample points (as compared to a low-order least-squares polynomial), but the resulting surface fits are generally more reliable than a polynomial surface fit. The better performance is a result of the incredible ability of a kriging model to fit to amorphous design contours that simply cannot be approximated well by a polynomial response. In short, if the user is able to afford the requisite number of points to generate a kriging space mapping relationship then the resulting surrogate model is very likely to be more accurate than any polynomial counterparts.

Surrogate Construction Through Space Mapping

Assuming the variable relationship and residual difference approximation are accurate over the region in which they were derived, the low-fidelity model in concert with the space mapping can be used as a surrogate model for the high-fidelity response. While the process of space mapping requires q high-fidelity analysis calls and a greater number of low-fidelity analysis calls, the end result can be applied as a surrogate to approximate the high-fidelity response at the computational

expense of the low-fidelity analysis. This provides the potential for optimization and design space explorations at a higher level of fidelity than was previously achievable due to the expense of the high-fidelity analysis.

The algorithm described in this chapter was developed in the hopes of leveraging information from a lower fidelity model to improve the accuracy of a surrogate model based on higher fidelity data. Chapter IV presents the application of this algorithm to synthesized high and low-fidelity model-pairings. Chapter V presents the application of this algorithm to a pair of analytical models in use in the aircraft design community at present. The resulting space-mapped surrogate models are then compared to surrogate models constructed using only the data from the high-fidelity model.

IV. Conceptual Applications

The first step of this research was the implementation of the theories outlined in Chapter III using analytical models defined specifically for this purpose. This step was necessary for the refinement of the algorithms used as well as an early investigation of the applicability of this method of surrogate construction. This chapter introduces the reader to the analytical models employed, the specific algorithm used, and the results of these conceptual applications.

High-fidelity and Low-fidelity Model Pairs

Three pairs of analytical models were formed for the purpose of demonstrating the algorithm detailed in Chapter III. Each pair consists of two models that serve as proxies for a high-fidelity and low-fidelity analytical tool set. It was important in the development and debugging of the theories and algorithms discussed in Chapter III for the analytical models to be computationally inexpensive yet retain an inherent relationship between the two models. The intention was for these conceptual applications to differ from the ESAV-relevant pairing used in Chapter V in computational expense only.

Case 1: 2-D Lift model vs. 3-D Lift model

The first model pair makes use of two related methods for the prediction of lift. The low-fidelity model is defined as the equation for predicting the lift of a two-dimensional airfoil.

$$R_L = \frac{1}{2} \rho V^2 S C_{L_\alpha} \alpha \quad (45)$$

The high-fidelity model is defined as the equation for predicting the lift of a three-dimensional airfoil.

$$R_H = \frac{1}{2} \rho V^2 S \left(\frac{C_{l_\alpha}}{1 + \frac{57.3 C_{l_\alpha}}{\pi e AR}} \right) \alpha \quad (46)$$

The low-fidelity analysis uses a lift-curve slope, C_{L_α} , equal to the two-dimensional lift-curve slope, C_{l_α} . In the high-fidelity analysis, the lift-curve slope is a function of the two-dimensional lift-curve slope, the efficiency factor of the wing, and the aspect ratio of the wing. The variable definitions for the high and low-fidelity models have been collected in Table 4.

Table 4. Variables present in the Case 1 analytical models

<i>symbol</i>	<i>definition</i>
shared variables	
ρ	density of air
V	velocity
S	planform area
C_{L_α}	lift-curve slope
α	angle of attack
additional variables	
e	efficiency factor
AR	aspect ratio

These two models were chosen because the relationship across fidelity levels is approximately linear with respect to angle of attack, illustrated in Figure 12.

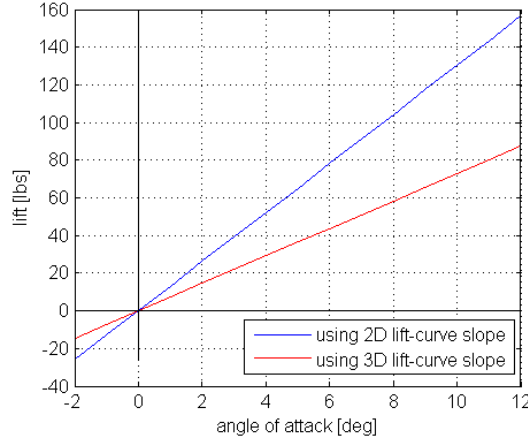


Figure 12. Difference in lift predictions by model (variables that are held constant are listed in Table 13 in the Appendix)

Case 2: 2-D Lift Model vs. 2-D “Saddle” Function

For the second pair of analytical models, the low-fidelity analysis proxy from the Case 1 model pairing was retained (Equation 45). For the high-fidelity analysis proxy, a modification to Equation 45 was crafted to introduce a nonlinear relationship between fidelity levels yet keep the same number of additional variables as Case 1. The additions to the 2-D lift equation are shown below:

$$R_H = \frac{1}{2} \rho V^2 S C_{L_\alpha} \alpha \sin(\bar{x}_{H_6}) ((\bar{x}_{H_7} - 1)^2 + 1). \quad (47)$$

Using Equation 47 as the high-fidelity proxy is beneficial because the relationship between fidelity levels changes depending upon the region of the design space. In a small region centered at values of 90 degrees for $\bar{x}_{H_6}$ and 1 for $\bar{x}_{H_7}$ the two analysis methods return approximately the same output, but for any other area of the design space the output values may vary drastically. The “saddle” title for this high-fidelity model refers to the appearance of the response surface when plotted against $\bar{x}_{H_6}$ and $\bar{x}_{H_7}$ (seen in Figure 23).

*Case 3: Truncated **peaks()** vs. **peaks()** Function*

The final pairing of analytical models include the **peaks()** function in MATLAB and a truncated version of the **peaks()** function for the high and low-fidelity models, respectively. These two models are useful in a conceptual application because the functions are dependent upon only two variables, which allows for a visual interpretation of the responses for the high-fidelity model, each surrogate model being evaluated, and the low-fidelity model. The equations for the **peaks()** function and the truncated **peaks()** function are shown below (the first term in Equation 48 is purposefully missing in Equation 49).

$$R_H = 3(1-x)^2 e^{(-x^2-(y+1)^2)} - 10\left(\frac{x}{5} - x^3 - y^5\right) e^{(-x^2-y^2)} - \frac{1}{3} e^{-(x+1)^2-y^2} \quad (48)$$

$$R_L = -10\left(\frac{x}{5} - x^3 - y^5\right) e^{(-x^2-y^2)} - \frac{1}{3} e^{-(x+1)^2-y^2} \quad (49)$$

The two models can be compared qualitatively through an inspection of their response surfaces plotted over the x and y variables (Figure 13). The low-fidelity model resembles its high-fidelity counterpart, but there is a single peak near the origin absent in the low-fidelity response. The truncation of the **peaks()** function is supposed to represent a scenario where the low-fidelity model is capable of providing a reasonable estimate of a real-world response over the majority of the design space. The added fidelity represented by the additional peak in the high-fidelity model is symbolic of the model’s ability to capture a high-order physical phenomenon absent in the low-fidelity model.

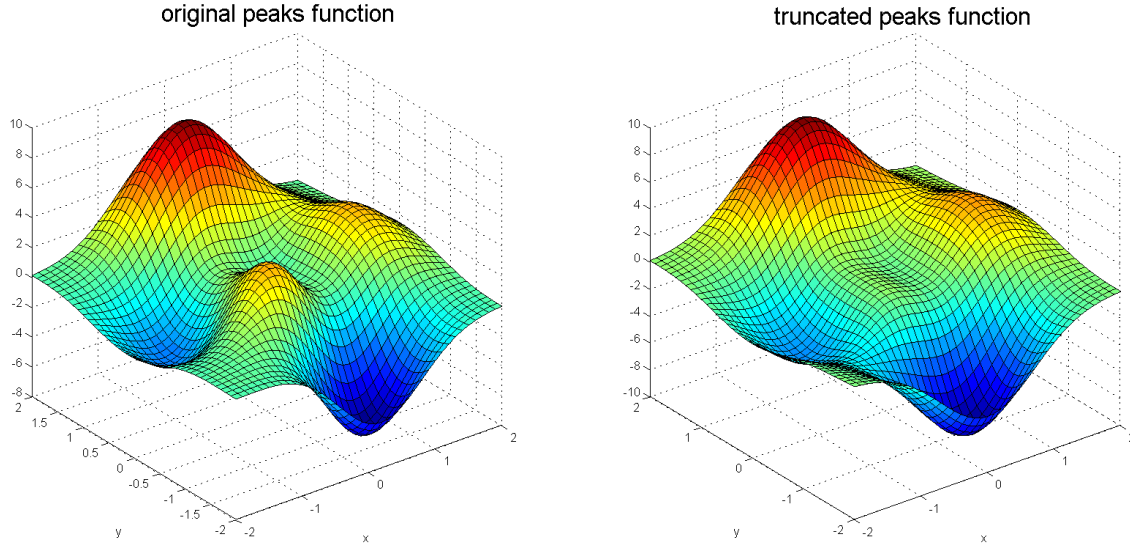


Figure 13. Visual comparison of the `peaks()` function versus the truncated `peaks()` function

Surrogate Model Comparison

The space mapping theory detailed in Chapter III is implemented on each of the three model-pairings just discussed. The following sections evaluate the performance of the resulting space-mapped surrogates in terms of the total error calculation (Equation 40) as well as a comparison of the high-fidelity and space-mapped surrogate response surfaces generated over the region in which $\mathbf{P}$ and $\Delta\mathbf{P}$ were derived. A qualitative comparison is generated by plotting the high-fidelity response against two of the design variables; the remaining design variables are held at constant values for plotting purposes. The associated quantitative comparison is the root mean square error (RMSE) between the high-fidelity response and the space-mapped surrogate response. The RMSE is calculated by taking the squared difference between the high-fidelity response and the space-mapped surrogate response at each intersection in the plotting grid and then summing each of these errors. The plotting grid is held at a constant size for each model pairing case so a direct comparison can be made.

The choice of what specific optimization scheme to use within the space mapping algorithm is left to the user. The optimizer used in Task 2 for the model-pairings is the unconstrained minimization function `fminunc()` included in the optimization software package for MATLAB . The options available within `fminunc()` were not varied to reduce the number of low-fidelity analysis calls required in the minimization sequence. Unless otherwise stated, the specific optimization

algorithm within `fminunc()` was set to sequential quadratic programming (SQP), and the tolerance for the change in value of the objective function was set to 1×10^{-5} for all cases. It is entirely possible the number of calls to the low-fidelity analysis could be reduced through a more tailored approach to the minimization sequence, but this was not attempted in this analysis. For each case, and for each space mapping form, the number of high and low-fidelity analysis calls will be documented.

In the fitting of a nonlinear variable relationship described in Task 3B, there is an additional optimization process by which the parameters of $\mathbf{P}$ and $\Delta\mathbf{P}$ are determined. The selection of an optimizer for this process is also left to user; for the conceptual applications shown below the genetic algorithm toolbox included in the optimization software package for MATLAB was used. Due to the standardization process applied to the design variables, approximately half of the scaled high-fidelity variable values will be negative. Initial attempts to raise the scaled high-fidelity design variables with negative values to non-integer powers resulted in complex numbers. These complex numbers were not useful in the context of the algorithm presented in Chapter III. A genetic algorithm has the ability to handle integer quantities, which proved useful in the selection of the powers ($\bar{\beta}$) listed in Equation 41 for negative variable values. The upper boundary for the element value of $\bar{\beta}$ is set to the appropriate degree of least-squares polynomial (for the number of sampled points, q) + 1, while the lower boundary is always set to 1.

Case 1 Space-Mapped Surrogate Models

Tasks 1 and 2 in the modified space mapping algorithm do not vary in implementation for the various assumed variable relationships, except for the number of sampled points. The upper and lower variable boundaries set in Task 1 for the sampling process are listed in Table 15 in the Appendix for each case. The first variable relationship assumed for the three model-pairings is the least-squares k^{th} polynomial. This process begins with an investigation into the number of high-fidelity responses required for different degrees of polynomials. The high-fidelity model has 7 design variables, and the least-squares fitting of coefficients requires at least 8 high-fidelity responses to create an over-determined system. In practice, however, the number of sampled high-fidelity data points (sampled by the latin hypercube spacing algorithm) was set to 12. More data points are recommended due to the erratic behavior seen in the response surfaces generated by 8 data points, as seen in the following figure. The surrogate models shown for the Case 1 model

pairing are plotted against the efficiency factor and the aspect ratio design variables. The high-fidelity response is the blue surface, and is generated by executing the high-fidelity function over the established ranges for efficiency and aspect ratio while holding the five shared design variables constant at the values listed in Table 14 in the Appendix. The gray surface is the surrogate generated using the low-fidelity analysis in conjunction with the space mapping as shown in Figure 4. The values for the variables held constant in the production of the response surfaces shown in subsequent figures are in Table 14 in the Appendix.

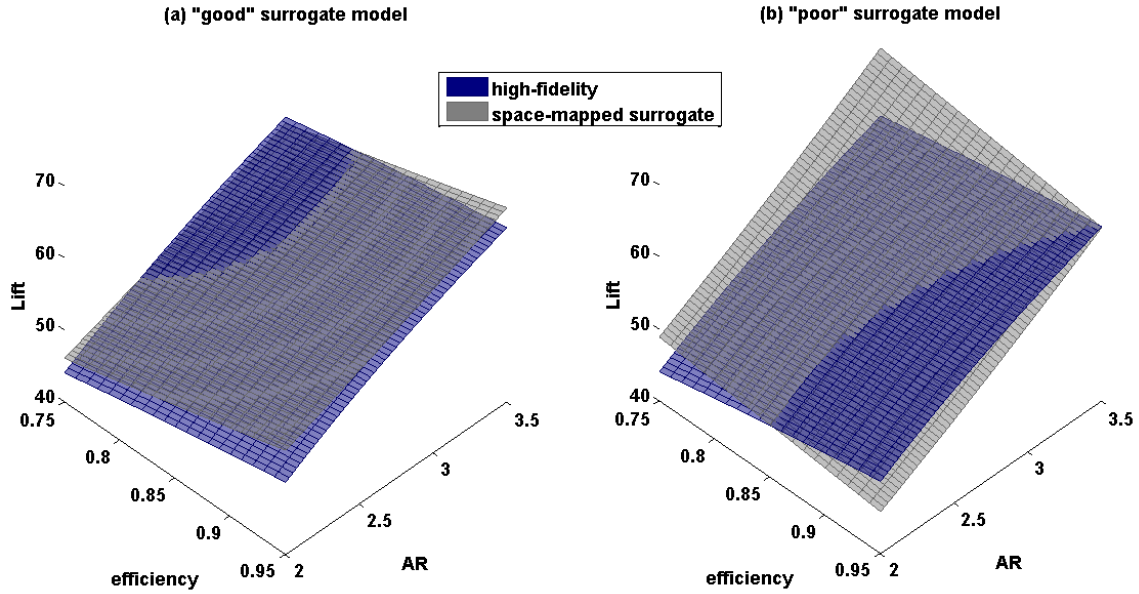


Figure 14. Visual comparison between two surrogate models with $q = 8$ where (a) is considered a “good” fit while (b) is considered a “poor” fit

The results were more repeatable when the number of sampled points was increased to 12. The instability of the resulting space-mapped surrogate behavior is a result of the random sampling process, which does not capture the necessary information from the higher fidelity level with the minimum number of samples from the design space. Increasing the number of sampled points increases the odds the sampled points will impart the relevant information to the lower fidelity design space. An example of a first-order polynomial space-mapped surrogate model using 12 sampled points is shown in Figure 15.

The first-order polynomial space-mapped surrogate captures the overall trends of the higher design space rather well. A first-order space mapping performs well in this context because the effects of the two additional design variables, efficiency factor and aspect ratio, are approximately

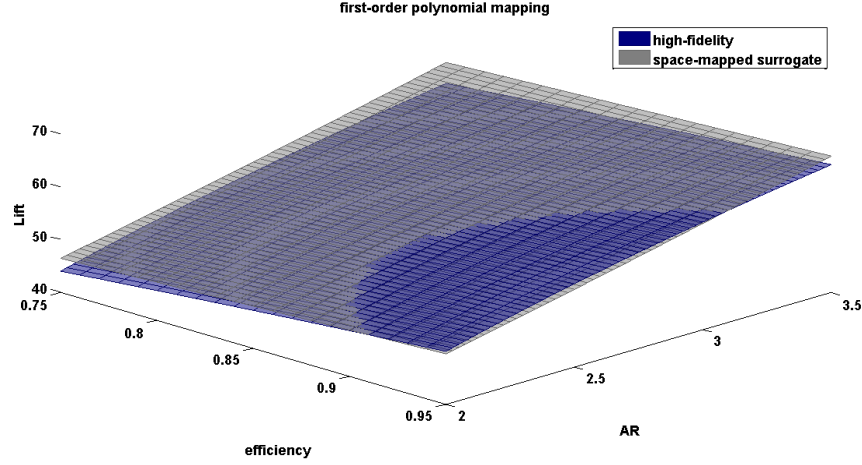


Figure 15. First-order polynomial space-mapped surrogate model using 12 sampled points

linear over the range of variable values under consideration. For model-pairings where the additional design variables do not affect the design space in a linear manner, a first-order space mapping will not likely yield a sufficiently accurate surrogate model. The number of high and low-fidelity analysis executions associated with the space mappings shown thus far are shown in the following table.

Table 5. Number of high and low-fidelity analysis calls for the first-order space mappings of the Case 1 model pairing

q	high-fidelity executions	low-fidelity executions	Figure
8	8	330	Figure 14a
8	8	330	Figure 14b
12	12	486	Figure 15

Next, a second-order polynomial is fit to the space mapping data and compared with a first-order polynomial fit using the total error calculation. Extrapolating the number of required sampled points from the lessons learned in the first-order case, the number of sampled points is set to 24 for this space mapping. The computed total error for the first-order and second-order polynomial space mappings are 15.608 and 7.965, respectively. The RMSE generated by the two surrogates shown in Figure 16 are 1,723.5 for the first-order polynomial and 646.7 for the second-order polynomial. These error calculations imply the second-order polynomial yields a more

accurate approximation of the high-fidelity response. This additional accuracy is gained at the cost of executing the high-fidelity analysis 24 times, and the low-fidelity analysis 984 times.

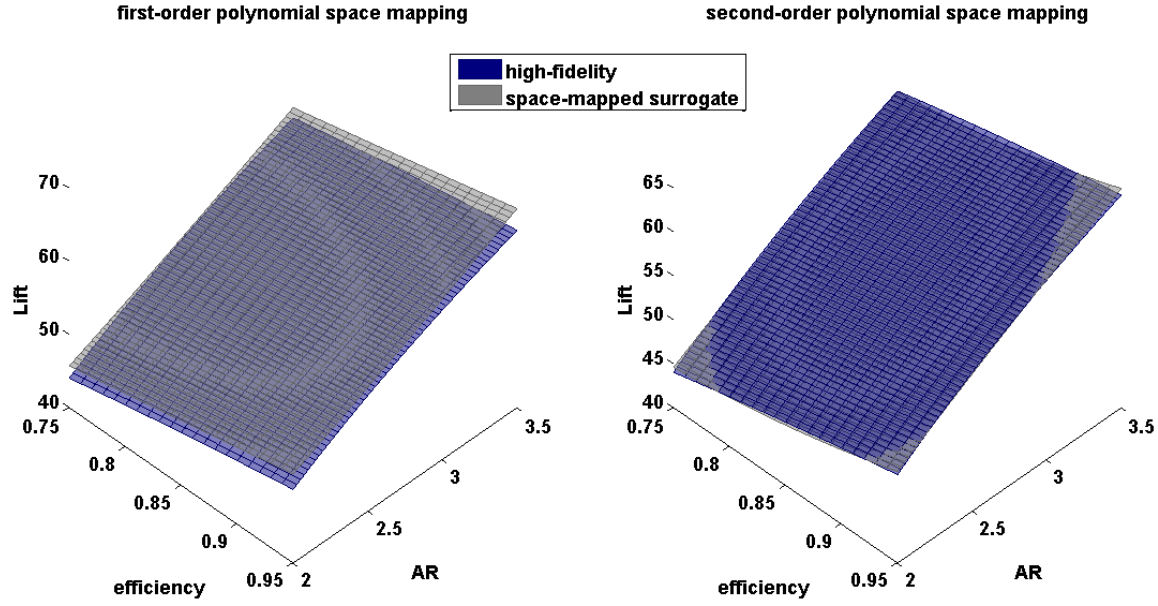


Figure 16. Visual comparison between surrogate models using a first-order and second-order least-squares space mapping

For the Case 1 model pairing, the total error for a given polynomial degree decreased as the degree of the polynomial form increased. The largest decrease in total error follows the transition between a first-order and a second-order polynomial form, with smaller decreases for each additional increase in k . The number of sampled points for the Case 1 model-pairing is set to 12 times the polynomial degree, which means for the data presented in Figure 17 each error calculation is the sum of all errors for 120 sampled points. The tenth-degree polynomial space mapped surrogate response is included in the Appendices (Figure 49). The number of low-fidelity analysis calls for $q = 120$ was 4,896.

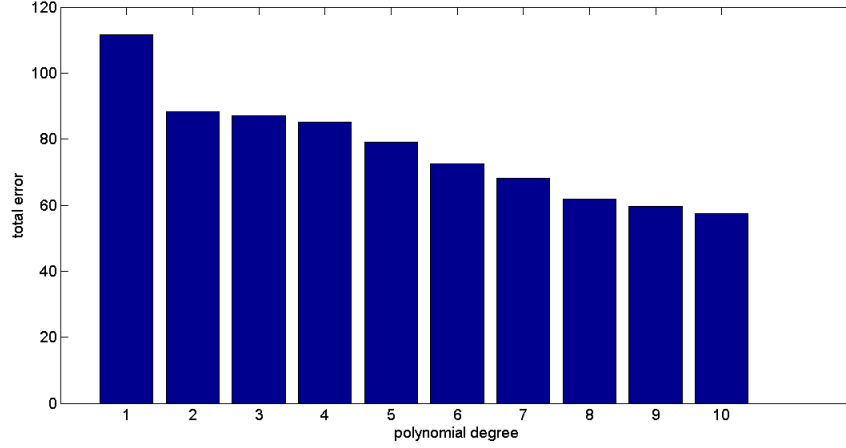


Figure 17. Total error calculations for each degree of polynomial fit to the space mapping data, $q = 120$

If the total error calculation is divided by the number of sampled points, the average error per sampled point can be used to compare space mappings using different numbers of sampled high-fidelity data points. Figure 18 shows this comparison metric for the Case 1 model pairing. The data in the figure were taken from single executions of the space mapping algorithm. Since the sampled high-fidelity design vectors are chosen through a latin-hypercube sampling method, there will be some variation in the error values for repeated executions of the space mapping algorithm. The data in the chart are not averaged values over multiple space mapping runs; the figure is intended to convey the trend that the error per point generally decreases as the degree of the polynomial increases.

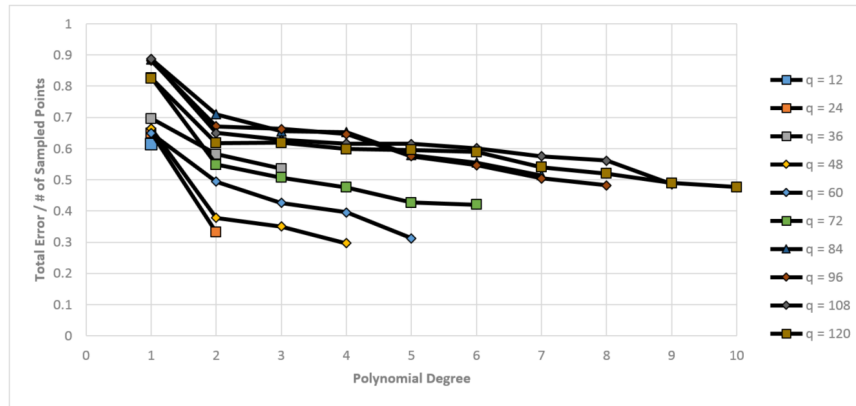


Figure 18. Average Errors per sampled point for a range of q showing a general decrease in error per point as the polynomial degree is increased

Using the same sampled data points as the surrogate model depicted in Figure 15 (where $q = 12$) and executing Task 3B results in the surrogate model shown in Figure 19. The maximum value for any element in the powers vector ($\bar{\beta}$) is set to 2, which is higher than the appropriate value of k for a least-squares polynomial fitting with only 12 sampled points. The nonlinear form returned is predominantly first-order due to the approximately linear response of the high-fidelity response itself. As a result, the surrogate models from Tasks 3A and 3B are very similar in appearance. The RMSE for the two responses are 2518.1 and 5201.7 for the least-squares polynomial and the nonlinear polynomial form, respectively. These RMSE values imply that the least-squares space mapping provides the better surrogate model for this model pairing and number of sampled points.

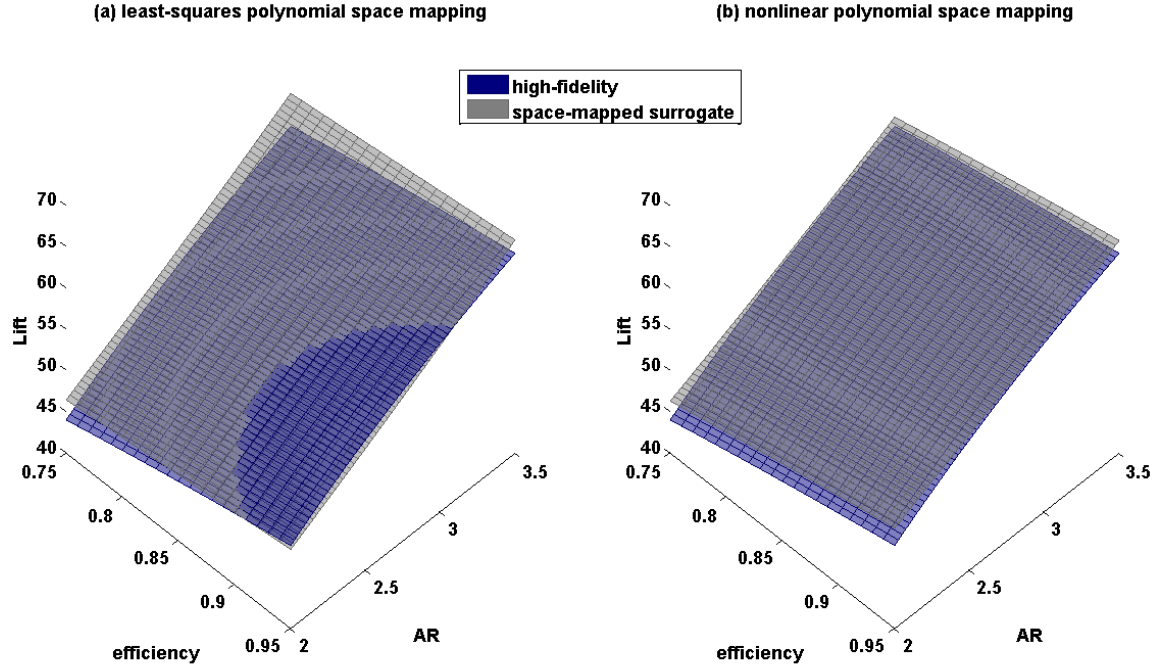


Figure 19. Comparison between the surrogate models constructed assuming (a) a least-squares polynomial and (b) a nonlinear space mapping for $q = 12$

Applying the steps in Task 3B to the same sampled points as for the surrogate model in Figure 16 ($q = 24$) and setting the maximum value for any element in $\bar{\beta}$ to 3 results in the surrogate model shown in Figure 20b. The process in Task 3B found the best variable relationship for many of the low-fidelity variables to be a linear one (powers equal to one, rather than two). The resulting surrogate model seems to perform poorly compared to the least-squares surrogate shown in Figure 20a. This is interesting because the initial expectation for this nonlinear form was

for the resulting surrogate model to be as good if not better than the surrogate model from the least-squares polynomial. Upon closer inspection, the least-squares variable relationship is better equipped (for this specific application) to capture the variable relationship due to the presence of all available powers for each high-fidelity variable. The coefficients that deal with the offset and the multipliers for the first-order design variables are able to capture the approximately linear relationship between the two fidelity levels. The coefficients that multiply the second-order design variables are then able to make more minute corrections, resulting in the better performing surrogate model depicted in Figure 16 and Figure 20a.

The nonlinear form prescribed in Task 3B, however, is limited to a single coefficient and a single power for each high-fidelity design variable. This leads the genetic algorithm to choose a predominantly linear relationship because this is the best fit for the data in this limited polynomial form. It should be noted that the process outlined in Task 3B can be applied to any particular form for a variable relationship. A different form, or even multiple forms, can be implemented at the discretion of the user and compared using the space mapping process outlined in Task 3B. The RMSE errors for the surrogates shown in Figure 20 are 646.7 and 3027.7 for the second-order polynomial and the nonlinear polynomial form, respectively.

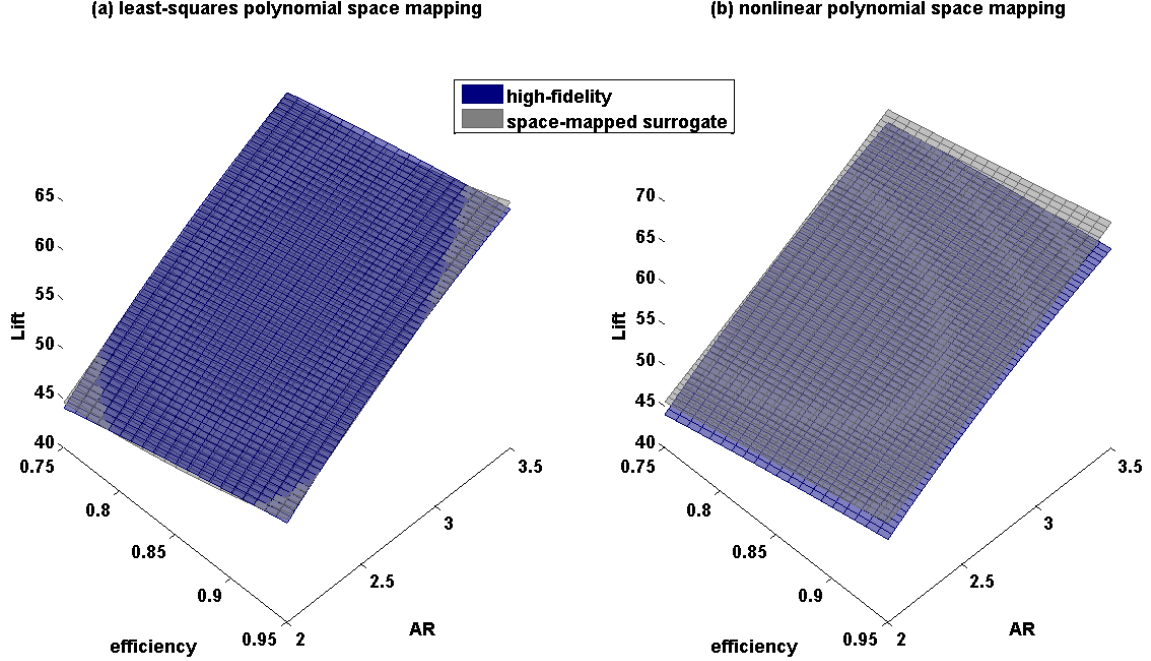


Figure 20. Comparison between the surrogate models constructed assuming (a) a least-squares polynomial and (b) a nonlinear space mapping for $q = 24$

The final space mapping relationship to be applied to the Case 1 model pairing is the nonlinear kriging form outlined in Task 3C. For a kriging model to be fit to the data collected in Tasks 1 and 2, a minimum of 36 high-fidelity data points are required (for reasons having to do with the rank requirements of a matrix in the kriging process). The number of sampled points is therefore set to 38 to be conservative and to allow for a more consistent space-mapped surrogate model. For a $q = 38$, the number of low-fidelity analysis calls was 1,560. As expected, the surrogate model built using the kriging relationship is the best performing surrogate model (from both a qualitative and quantitative standpoint). A third-degree polynomial is plotted for comparison ($k = 3$ is appropriate for the number of sampled points and this model pairing). The RMSE values for the surrogates shown in Figure 21 are 417.1 and 34.6 for the third-order least-squares polynomial and the kriging space mapping relationships, respectively.

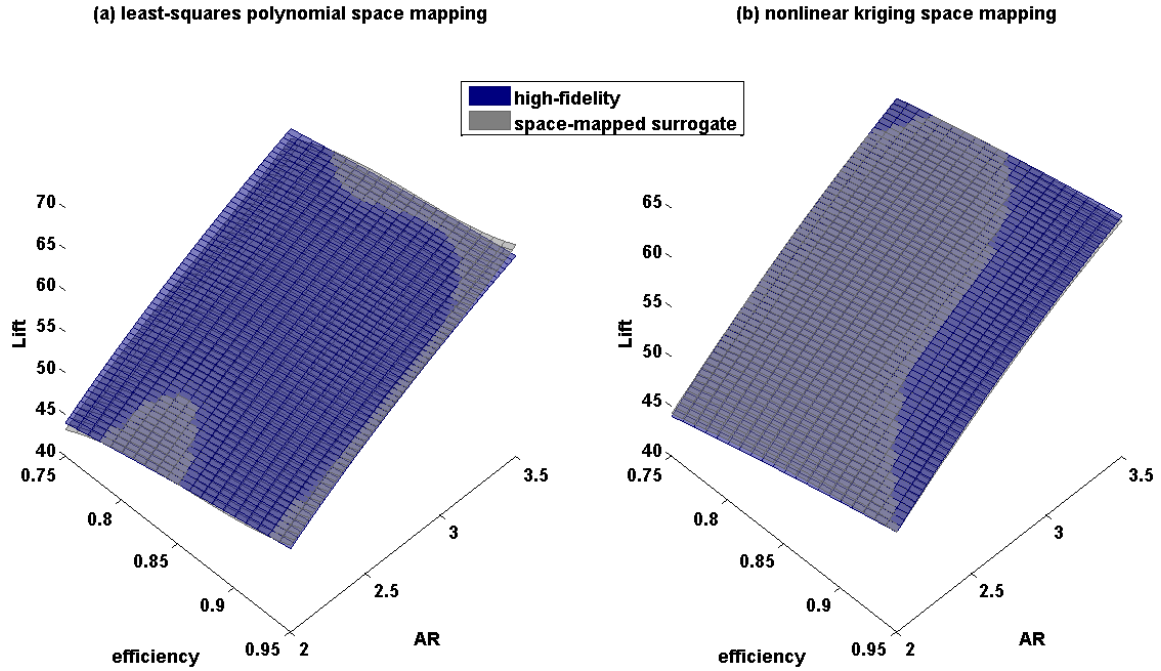


Figure 21. Comparison between the surrogate models constructed assuming (a) a least-squares polynomial and (b) a nonlinear space mapping using kriging models for $q = 38$

There is a problem of practicality associated with the kriging form used in Task 3C that deserves mentioning. The initiation of a kriging model requires a set number of sampled data points from the high-fidelity model, and this set number is typically larger than the number of samples required in either Tasks 3A or 3B. When executing Task 3C, the number of sample points required to fit kriging models to each low-fidelity design variable is the same number required to fit

a kriging model to the high-fidelity response. As such, it is only fair to compare the surrogate model derived from a kriging implementation of space-mapping with a traditional kriging surface constructed from the actual high-fidelity responses (R_H). For the Case 1 model pairing, the traditional kriging model performs better than the space-mapped surrogate when comparing the RMSE values (17.6 for the traditional kriging model as compared to 34.6 for the space-mapped surrogate). The qualitative comparison between the two response surfaces is shown in Figure 22.

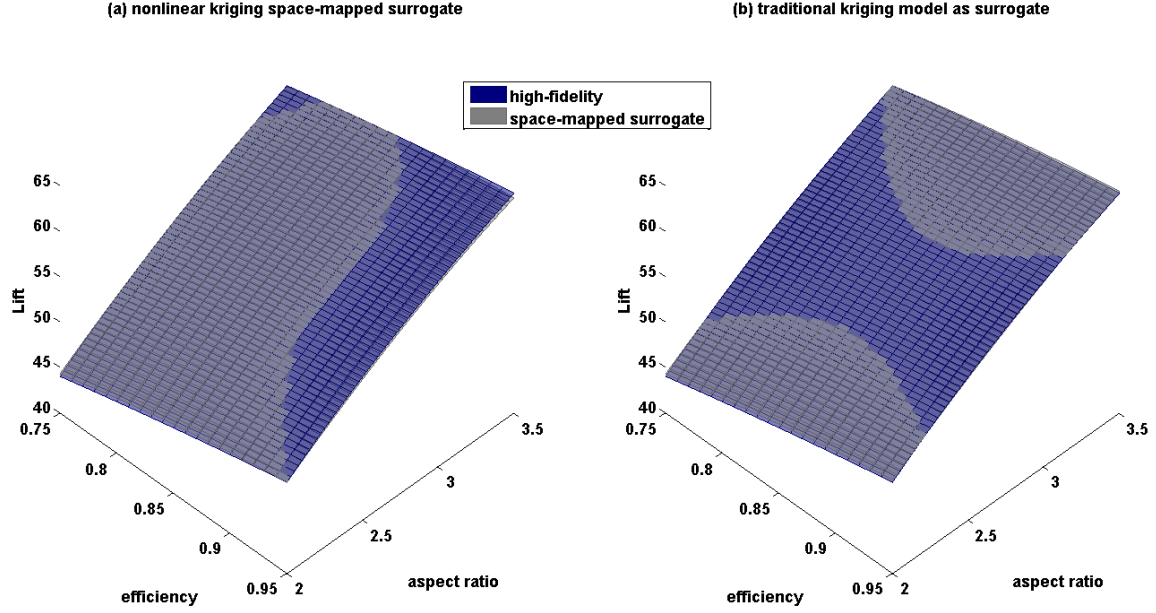


Figure 22. Comparison between the surrogate models constructed through (a) a kriging implementation of space mapping and (b) a traditional kriging model acting as a surrogate

Case 2 Space-Mapped Surrogate Models

While the first model pairing exhibits an approximately linear relationship between the fidelity levels, the second model pairing was engineered to exhibit a nonlinear relationship with respect to the additional variables in the high-fidelity model. The model pairing is nicknamed the “saddle” pairing because of the distinctive features of the high-fidelity response when plotted against the two additional design variables, $\bar{x}_{H_6}$ and $\bar{x}_{H_7}$. In the subsequent plots which compare the high-fidelity response to the space-mapped surrogate models, the high-fidelity response surface is calculated by holding the shared variables constant (see Table 14) while varying the two

additional design variable values over the ranges shown in the plots. As before, the blue surface depicts the high-fidelity response while the gray surface depicts the surrogate response.

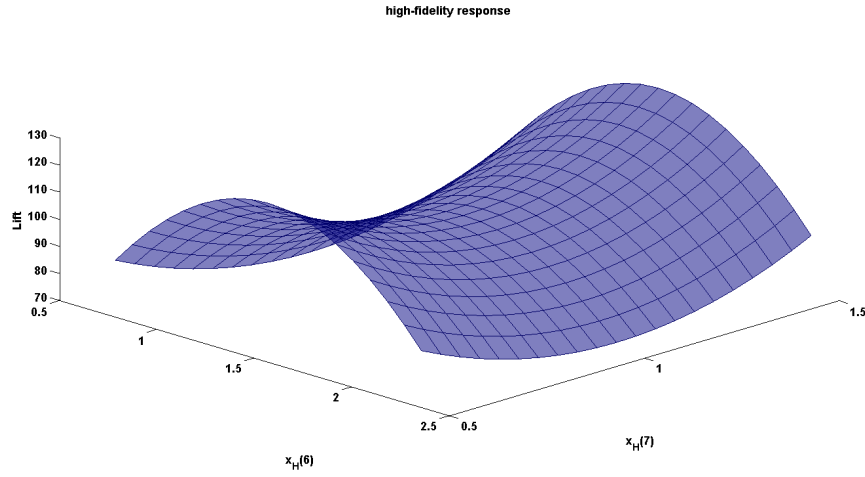


Figure 23. High-fidelity model response over a range of values for $\bar{x}_{H_6}$ and $\bar{x}_{H_7}$

The first space mapping explored for this model pairing is a linear least-squares assumption for $\mathbf{P}$ and $\Delta\mathbf{P}$. Based on the experience with the first model pairing (with an equal number of high and low-fidelity design variables), the number of sampled points is set to 12. The resulting surrogate model from this space mapping is shown in Figure 24. The surrogate response is a linear approximation of the high-fidelity behavior, and so it appears as a plane in 3D space. The orientation of the plane should be parallel with the $\bar{x}_{H_6}$ - $\bar{x}_{H_7}$ plane due to the symmetry of the high-fidelity response over the range of plotting variables, but is not due to slight biases in the sampled design vectors involved in Task 1. When the space mapping algorithm is repeated for different sampled points, the first-order surrogate response fluctuates about the point $[\bar{x}_{H_6} = \pi/2 \text{ radians}, \bar{x}_{H_7} = 1.0]$. What is clear from Figure 24 is that a linear form for $\mathbf{P}$ and $\Delta\mathbf{P}$ is insufficient to capture the nonlinear behavior absent in the low-fidelity model. For the 12 sampled high-fidelity executions in this space mapping, there were 384 low-fidelity executions performed.

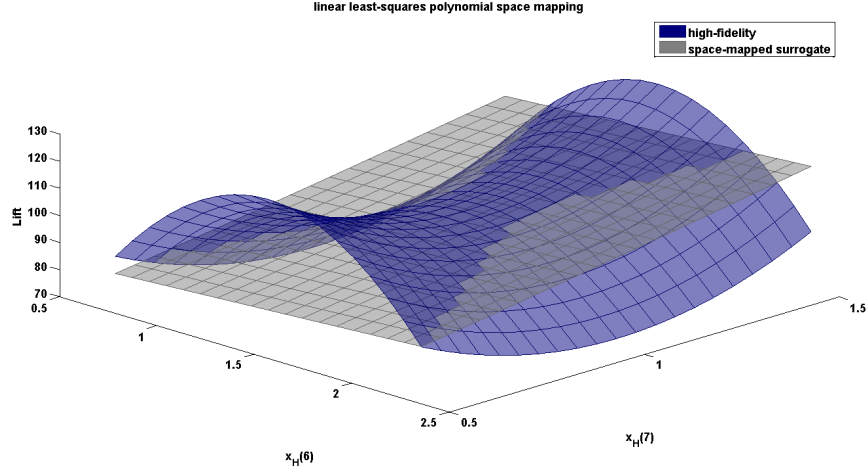


Figure 24. Surrogate model constructed using 12 sampled high-fidelity data points and a linear least-squares space mapping

For the nonlinear behavior present in the high-fidelity response to be imparted on the low-fidelity response, a nonlinear relationship for $\mathbf{P}$ and $\Delta\mathbf{P}$ is needed. Increasing the number of sampled points to 24 and allowing for a second-order least-squares polynomial fitting yields the space-mapped surrogate model shown in Figure 25.

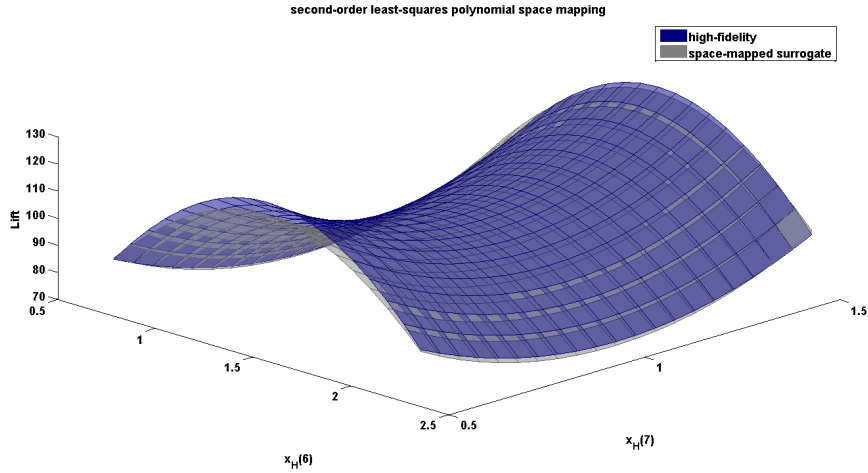


Figure 25. Surrogate model constructed using 24 sampled high-fidelity data points and a second-order least-squares space mapping

This second-order relationship exhibits the ability to capture the nonlinear behavior present in the high-fidelity response rather well. This model pairing (over the range of variables for which the space mapping was derived) is well represented by a second-order relationship between fidelity

levels, so the performance of the second-order least-squares polynomial is not very surprising. For the 24 sampled high-fidelity design points there were 762 low-fidelity analysis executions.

Subsequent increases in the order of the least-squares polynomial form provide marginally better total error values, with the major gain in surrogate model performance occurring in the transition from a linear variable relationship to a second-order one. A similar comparison to the one shown in Figure 18 for different degrees of least-squares polynomial forms for this model pairing is shown in Figure 26. Without any prior knowledge of the relationship between the two fidelity levels, a user could compare the average error per sampled point to determine the appropriate degree of polynomial (which will determine the number of points to sample in the high-fidelity design space).

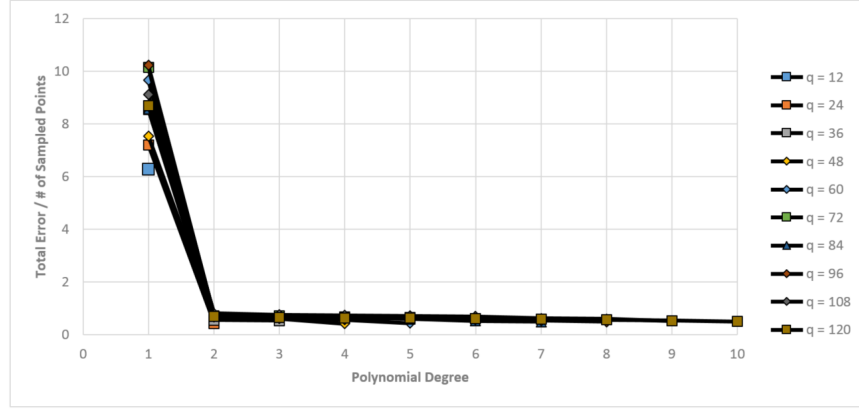


Figure 26. Average Errors per sampled point for a range of q 's showing a steep decrease from order 1 to 2, and marginal decreases for subsequent polynomial degrees

Figure 27 showcases the similarity between a second-order and third-order least-squares space-mapped surrogate models. There are no significant visual differences between the two models, which supports the conclusion drawn from the data in Figure 26 that polynomial degrees greater than 2 are not worth the extra high-fidelity model executions. The surrogate models constructed from fourth-order and higher least-squares polynomial forms were also visually indistinguishable from the second-order least-squares space-mapped surrogate model shown in Figure 25. The RMSE values for the two surrogate models are 872.8 for the second-order polynomial and 1115.9 for the third-order polynomial, implying that the second-order polynomial space mapping is the best polynomial representation of the high-fidelity response.

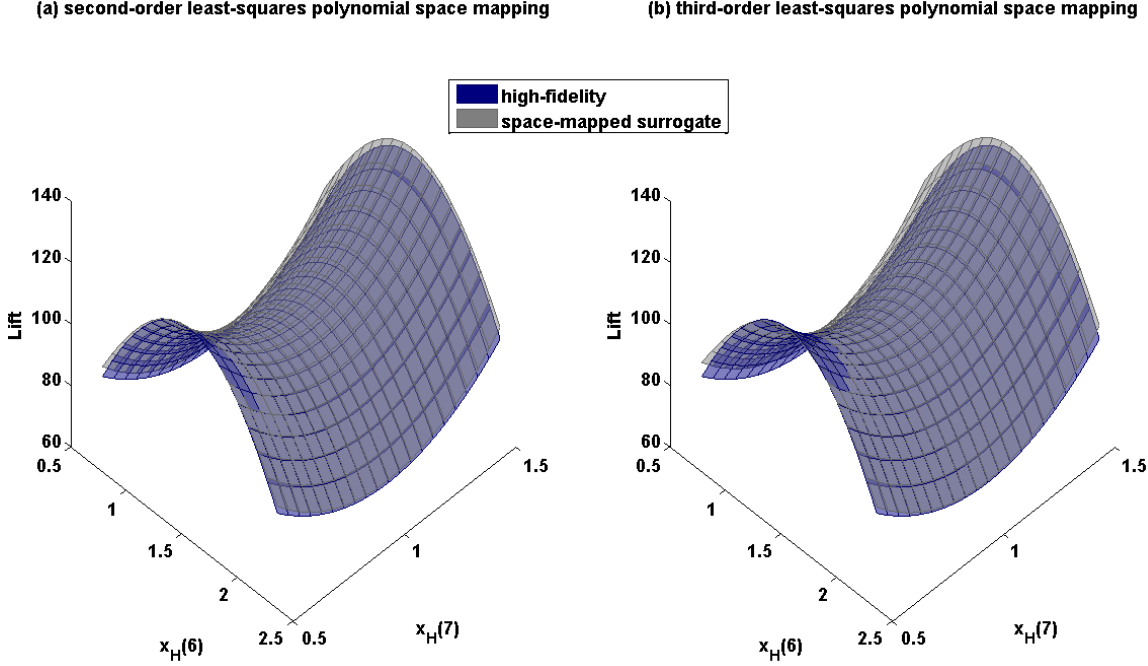


Figure 27. Surrogate models constructed using 36 sampled high-fidelity data points and assuming a (a) second-order and a (b) third-order least-squares space mapping

The assumption of the nonlinear form shown in Task 3B for $\mathbf{P}$ and $\Delta\mathbf{P}$ produces a surrogate model that actually outperforms the least-squares space mapping process in Task 3A with respect to replicating the nonlinear high-fidelity response for the least number of sampled points. For the same sampled high-fidelity design vectors as the first-order least-squares polynomial in Figure 24, the nonlinear form produces the surrogate model shown in Figure 28b. Using only the 12 sampled points, a nonlinear form of $\mathbf{P}$ and $\Delta\mathbf{P}$ were found to replicate the second-order behavior seen in the high-fidelity response. This is a huge advantage over the second-order least-squares method, but the advantage is specific to the model pairing in question. After all, the least-squares space mapping process was found to be the better performing method in the first model pairing.

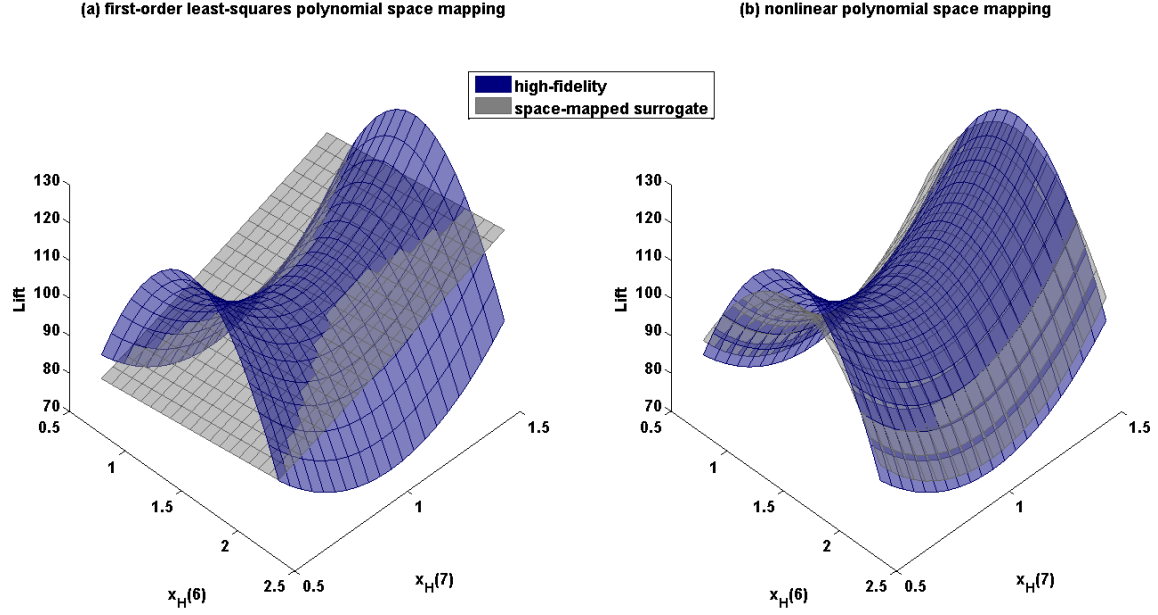


Figure 28. Surrogate models constructed using 12 sampled high-fidelity data points and assuming a (a) first-order least-squares and a (b) nonlinear polynomial space mapping

The performance comparison (in terms of average error per sample point) for the least-squares polynomial and the nonlinear polynomial forms in the space mapping process is shown in Figure 29. Note the low average error values for the second-order least-squares polynomial space mapping for sample points of 16 and greater. This data indicates that the number of samples taken in earlier space mappings ($q = 24$) may have been overly conservative for this model pairing.

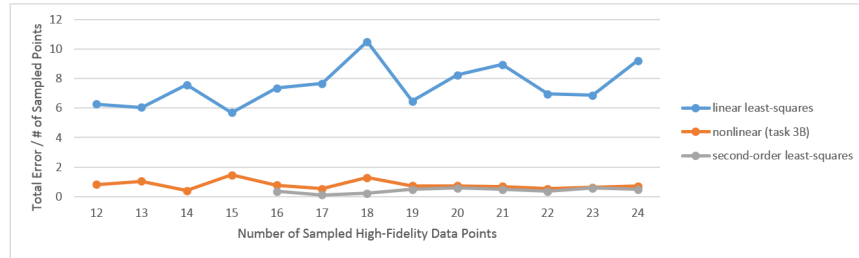


Figure 29. Comparison of average error per sampled point for a range of q using a first and second-order least-squares polynomial and a nonlinear polynomial space mapping approach

The final space mapping form applied to the second model pairing is the nonlinear kriging models for $\mathbf{P}$ and $\Delta\mathbf{P}$. The number of sampled points is set to 38, which is two more than the minimum number of responses for a kriging model to be built for the given number of high-fidelity design variables. As expected, the kriging implementation of space mapping is able to replicate the

high-fidelity behavior in the surrogate response. Figure 30 compares the space-mapped surrogate with a traditional kriging model constructed from the same 38 sampled data points. The space-mapped surrogate displays a better qualitative fit, but both surrogate models are able to recreate the trends seen in the high-fidelity design space. In quantitative terms, the RMSE for the kriging implementation of space mapping was 1334.0 while the RMSE for the traditional kriging model was 9434.9 (Figures 30a and 30b). For 38 sampled high-fidelity data points, the space mapping algorithm executed the low-fidelity analysis 1,218 times.

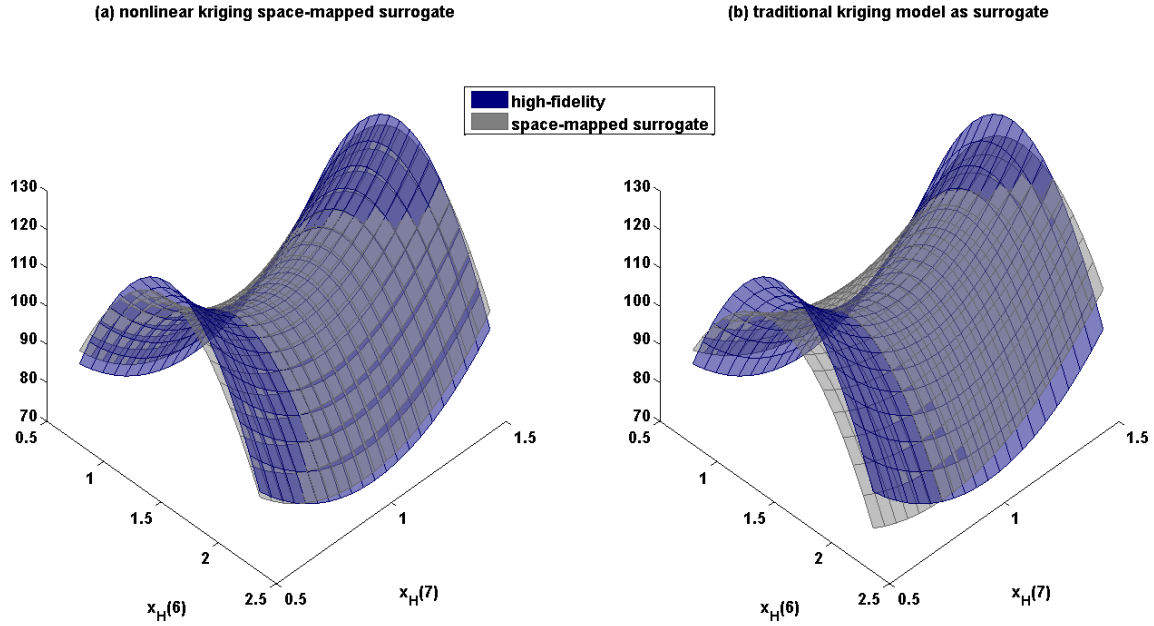


Figure 30. Surrogate models constructed using 38 sampled high-fidelity data points and assuming a (a) kriging implementation of space mapping and a (b) traditional kriging model acting as a surrogate

Case 3 Space-Mapped Surrogate Models

The Case 3 model pairing involving the `peaks()` function and the truncated `peaks()` function is certainly the most challenging of the three cases for the space mapping process. These functions are highly nonlinear, and have been used extensively in the field of gradient-based optimization to test the performance of optimization routines. The response surfaces shown in the following plots are not easily represented by polynomial expressions, and as such many of the space mapping forms detailed in this research struggle to replicate the high-fidelity response in the surrogate models. A kriging implementation of the space mapping process is somewhat successful, but no

more successful than the traditional kriging surface generated from the same sampled data points. As such, the Case 3 model pairing is presented to highlight potential limitations of the space mapping process to highly nonlinear applications.

The visual comparison between the `peaks()` function and the truncated `peaks()` function shown in Figure 13 is plotted over a wide range of x and y values to showcase the many peaks and valleys present in the design space. A subspace of the design space has been selected from which the sample points are taken; this subspace is illustrated in Figure 31. This subspace is focused on the area of the shared design space where the disagreement between the two fidelity levels is greatest. For each of the following space-mapped surrogates, 15 points were sampled from the high-fidelity subspace and the space mapping process required 405 low-fidelity analysis executions. Due to the nonlinearity of the surface being sampled, it should be noted that the surrogate models shown are a reflection of the specific points taken from the high-fidelity response. It is possible, with the appropriate selection of sample point locations, to achieve surrogate models that perform either very well or very poorly.

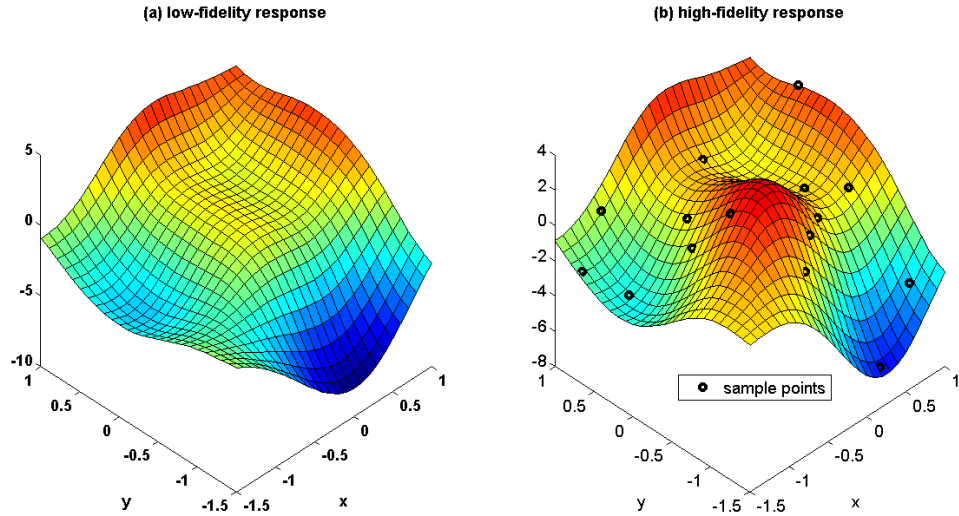


Figure 31. Illustration of the subspace, encompassing the peak in the high-fidelity model that is absent in the low-fidelity model, from which data points are sampled for the space mapping process

An application of the least-squares space mapping process of Task 3A yields a number of possible surrogates based upon varying degrees of least-squares polynomials. For 15 sampled data points, the maximum polynomial degree is set to 5 (a first-order least-squares fit requires 3 sample points). The polynomial degree with the least total error is the highest polynomial form applied,

and the resulting surrogate model is shown in Figure 32. While regions at the boundaries of the subspace are altered by the space mapping process, the surrogate model is able to replicate aspects of the peak from the high-fidelity response. The progression of surrogate models from a linear least-squares to the fifth-order least-squares space mapping is interesting because the surrogate’s ability to model the missing peak increases with each increase in polynomial degree. This trend ends when a sixth-degree polynomial is fit to the space mapping data; For those interested, the surrogate models based on the least-squares polynomials mentioned here are included in the Appendix in Figures 50 - 54.

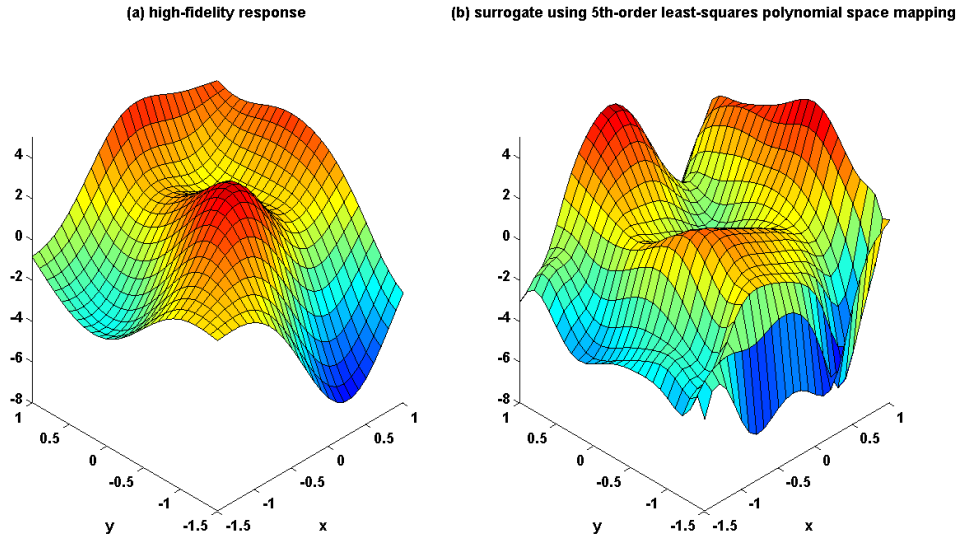


Figure 32. Surrogate model constructed using a fifth-order least-squares space mapping approach

The nonlinear polynomial form for the space mapping data prescribed in Task 3B is not so capable as the least-squares polynomial for this model pairing. The resulting surrogate model modifies the low-fidelity design space by forming a “plateau” in the region where the peak is present in the high-fidelity design space. This surrogate model may be more representative of the high-fidelity design space than the original low-fidelity model, but there is not a local maximum present in the surrogate design space in the vicinity of the high-fidelity peak. An optimization process applied to this surrogate model would therefore not yield the same local maxima and minima as the high-fidelity model, which is a desirable property for a surrogate model. This surrogate model is plotted alongside the high-fidelity design space in Figure 33. The RMSE for this surrogate model is 1890.5, which is lower than the RMSE value of the least-squares polynomial. The RMSE is 2223.6 for the fifth-degree polynomial space-mapped surrogate shown in Figure 33.

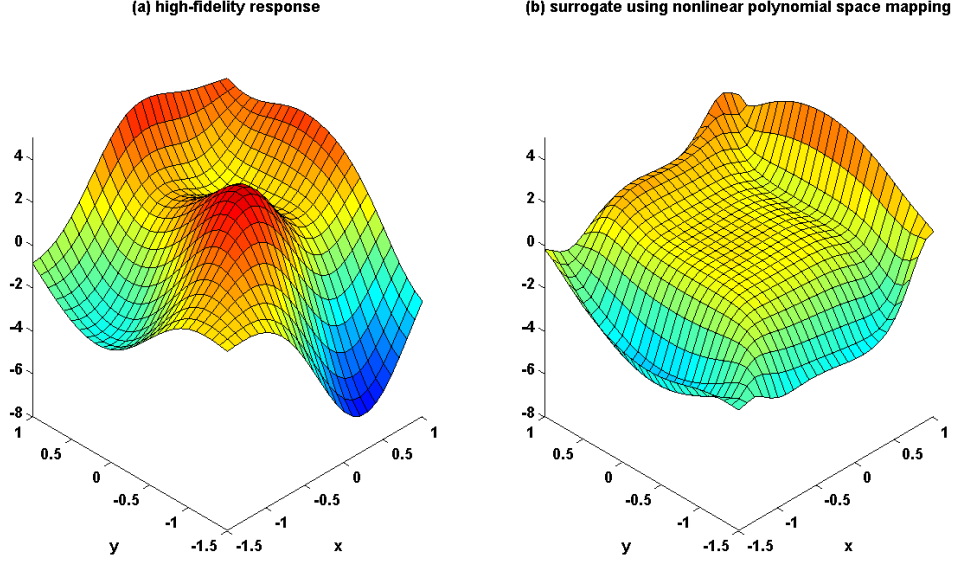


Figure 33. Surrogate model constructed using a nonlinear polynomial space mapping approach

Executing Task 3C involves fitting a kriging model to each low-fidelity design variable as well as the residual differences between fidelity levels. The resulting surrogate models constructed using this type of space mapping tend to reflect the presence of the peak in the high-fidelity model, but when compared to the traditional kriging response surface, constructed from the sampled high-fidelity data points, the surrogate model response surfaces tend to be much more jagged. For the 15 sampled data points shown in Figure 31, the surrogate model represents the high-fidelity response reasonably well (shown in Figure 34). The traditional kriging surface illustrated in Figure 35 arguably performs better due to the lack of the artificial ridges on the peak's surface. The performance of the kriging space-mapped surrogate model as well as the traditional kriging surface is heavily dependent upon the location of the sampled high-fidelity data points, and so a number of alternate sample datasets were run through the space mapping algorithm. The comparisons between these surrogate models and their traditional kriging surface counterparts are included in the Appendix in Figures 55 - 57. In general, for the many different sets of sampled points explored in this research for this model pairing, the traditional kriging surface was judged to be the better surrogate for the high-fidelity model. The RMSE for the kriging implementation of space mapping is 808.7 (response shown in Figure 34) while the RMSE for the traditional kriging model is 531.7 (Figure 35).

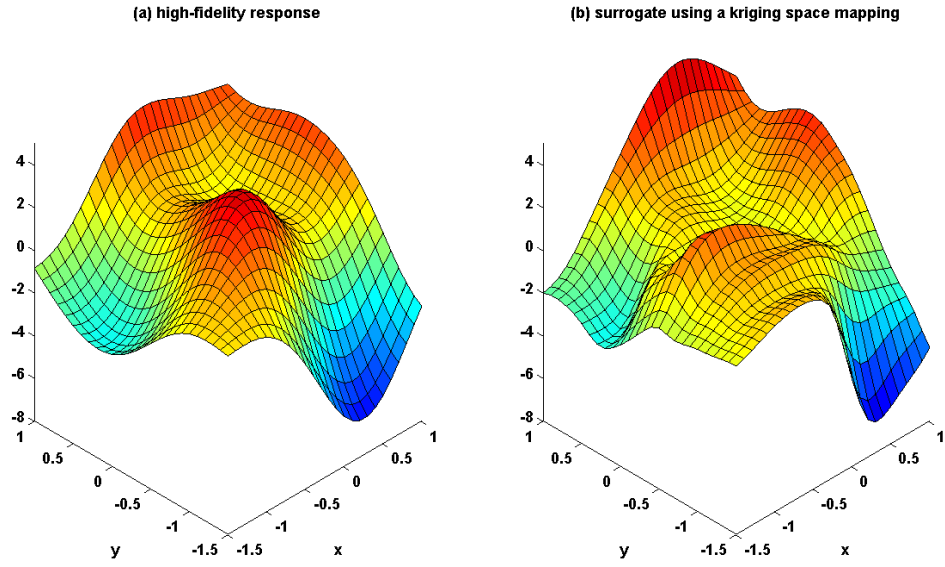


Figure 34. Surrogate model constructed using a nonlinear kriging space mapping approach

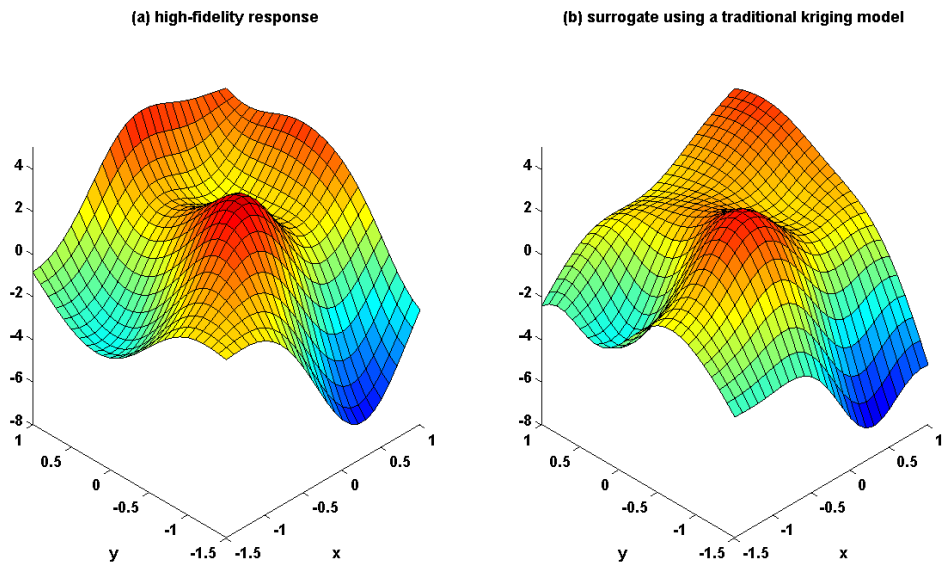


Figure 35. Surrogate model constructed using a traditional kriging response surface

V. Space Mapping Application with ESAV Tools

The Aerospace Vehicles Directorate of the Air Force Research Laboratory (AFRL/RQ) is interested in the optimal design of an Efficient Supersonic Air Vehicle (ESAV). In pursuit of this goal, the Multidisciplinary Science and Technology Center (MSTC) within AFRL/RQ is investing resources towards a design and optimization framework that will allow designers to capture the effects of innovative design features (such as advanced engine technologies, active aero-elastic wing, gust load alleviation, and tailless supersonic flight) early on in the design cycle [12]. The ability to capture subsystem interactions at a high level of fidelity drives the resulting design framework towards computationally expensive analysis methods.

As the computational cost of executing a single analysis increases, the associated cost of optimization using this analysis increases as well. As the analyses within the design framework are pushed to higher fidelity levels, the resulting time and resources required to execute a design iteration makes an optimization scheme infeasible. Space mapping may allow optimization at these higher fidelity levels using lower fidelity codes at the expense of the space mapping process outlined in Chapter III. The design framework used in this research is the ESAV model built within the Service-Oriented Computing Environment (SORCER) developed internally at AFRL/RQ. This framework contains analysis blocks representing the relevant engineering disciplines involved in the design of an efficient supersonic air vehicle.

The various analysis blocks are organized within the SORCER environment according to the N^2 diagram shown in Figure 36. The framework for the ESAV design begins with an analysis block that takes in inputs related to the size and shape of the vehicle and outputs various geometry files needed in later blocks. This analysis is performed by a tool called MSTCGeom which was developed in-house at RQ for ESAV design purposes. The MSTCGeom tool forms a finite-element model (FEM) for a conceptual vehicle, which is necessary for the various analysis blocks further down the N^2 flow. The automation of the FEM synthesis in the MSTCGeom tool greatly increases the utility of the ESAV analysis as a whole because manual FEM construction is a tedious and time-consuming task that can limit the number of design concepts analyzed in a given phase of development. The automation of the FEM construction allows a greater number of design configurations to be analyzed given time constraints on the phase of the design cycle [13].

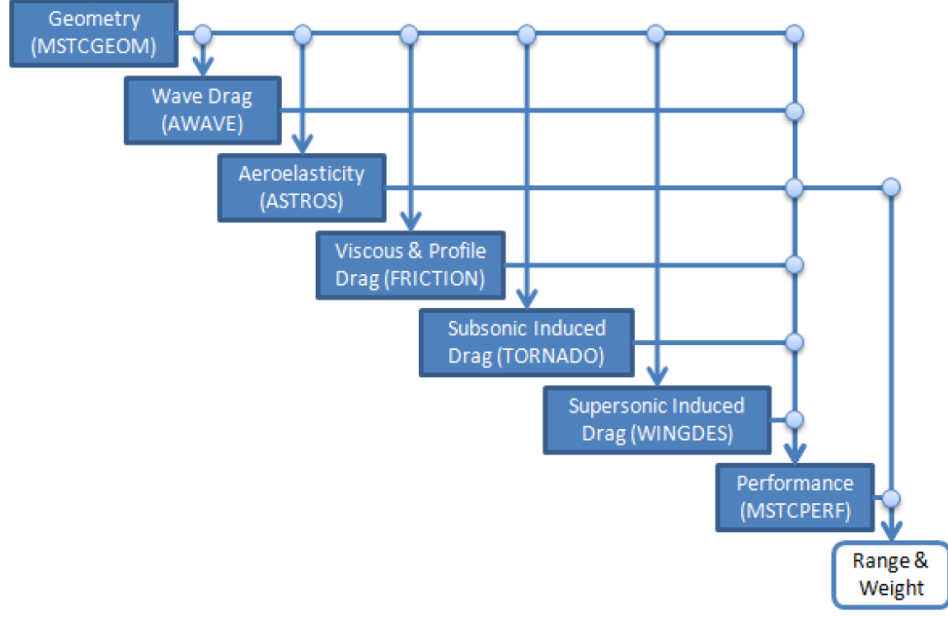


Figure 36. N^2 diagram depicting the various analysis blocks in the ESAV model within SORCER, taken from Ref [14]

ESAV Space Mapping Implementation

A small subset of the tools available within the ESAV design framework are analyzed in this research. The space mapping process outlined in Chapter III is applied to construct and evaluate a surrogate model for the ASTROS structural sizing tool. ASTROS is a tool developed for AFRL that is capable of analyzing and optimizing aerodynamic structures considering the various static and dynamic loads involved in the flight profile. ASTROS takes a set flight profile as well as a FEM model of the aerospace structure and outputs (among numerous other things) the weight of the structure optimized for minimum mass, while withstanding the applied loads [15]. The FEM is constructed by the MSTCGeom tool, which requires the input variables listed in Table 6. The inputs to the MSTCGeom tool are the high-fidelity design variables for the ESAV space mapping application, and the ASTROS weight value is the high-fidelity response that is approximated.

Table 6. High and low-fidelity design vector for the ESAV space mapping application

high-fidelity design vector, $\bar{x}_H$		
<i>variable</i>	<i>definition</i>	<i>units</i>
area	planform area of the wing	in ²
aspect ratio	span ² /area	-
t/c @ root	thickness to chord length ratio at wing root	-
t/c @ tip	thickness to chord length ratio at wing tip	-
wing sweep	sweep of the wing (at the quarter chord position)	degrees
taper ratio	ratio of the wing's tip length to root length	-
camber location	location of the camber for the airfoil cross-section	%
max camber location	location of the maximum camber for the airfoil cross-section	%
wing twist	angle of twist at the wing tip of the airfoil cross-section	degrees
low-fidelity design vector, $\bar{x}_L$		
area	trapezoidal wing area	ft ²
aspect ratio	span ² /area	-
t/c @ root	thickness to chord length ratio at wing root	-
wing sweep	sweep of the wing (at the quarter chord position)	degrees
taper ratio	ratio of the wing's tip length to root length	-

The low-fidelity analysis for this ESAV space mapping is a weight prediction model based upon empirical data coupled with sizing approximations for the various subcomponents of an aircraft. This model is presented in [16] on pages 583-595. The number of variables input into this weight predictor is immense, but most of these inputs are held constant for the purposes of this space mapping. A full listing of these input parameters (and their values) is listed in the Appendix in Table 17.

The ESAV design framework used in this space mapping application is set to replicate the weight and planform characteristics of an F-16 aircraft (used within RQ for code validation purposes). All input parameters to the ESAV model not pertaining to the wing of the vehicle have been set to the values for an F-16. The same is true of the weight predictor, where the various inputs listed in Table 17 are set to values representative of an F-16. The design variables for both the high and low-fidelity models describe the wing structure for the vehicle. Table 6 lists the variables in both the high and low-fidelity design vectors, along with their definitions. One of the assumptions in this space mapping algorithm is the two models, while at differing levels of fidelity, share some commonality in the contours of the shared design space. For the ESAV space mapping application, there are five shared design variables and four additional high-fidelity design variables,

as shown in Table 7. Any improvements in predictive accuracy of the space-mapped surrogate models hinges upon the validity of this assumption.

Table 7. Shared-fidelity design vector for the ESAV space mapping application

shared-fidelity design vector, $\tilde{x}_L$		
<i>variable</i>	<i>definition</i>	<i>units</i>
area	trapezoidal wing area	ft ²
aspect ratio	span ² /area	-
t/c @ root	average thickness to chord ratio for the main wing	-
wing sweep	sweep of the wing (at the quarter chord position)	degrees
taper ratio	ratio of the wing's tip length to root length	-
additional high-fidelity design variables, $\tilde{x}_H$		
t/c @ tip	thickness to chord length ratio at wing tip	-
camber location	location of the camber for the airfoil cross-section	%
max camber location	location of the maximum camber for the airfoil cross-section	%
wing twist	angle of twist at the wing tip of the airfoil cross-section	degrees

In the application of the space mapping process, some of the shared design variables need to be converted from the format expected of the high-fidelity model to the format required in the low-fidelity model. The area variable needs minor alterations when passing from the ASTROS model to the weight predictor, and vice versa. For instance, the weight predictor considers the planform area of the wing to include the sections enclosed in the fuselage of the vehicle. ASTROS, on the other hand, only considers the wetted-area of the wing and disregards the wing structure within the bounds of the fuselage section. Additionally, the ESAV model inputs are for the vehicle's half-span while the weight predictor inputs are the vehicle's full wing-span. Lastly, the area variable is converted as appropriate between square inches and square feet. All of these conversions are needed in order to translate a high-fidelity design vector, $\tilde{x}_H$, into space-mapped low-fidelity design variables.

For the implementation of the space mapping algorithm, a number of datasets were extracted from the ESAV design framework using the latin hypercube sampling process. Since there are nine high-fidelity design variables, the size of the sample datasets were set to multiples of 14 (this was judged to be a conservative number of points necessary for implementing a linear least-squares space mapping). Three datasets were formed with 14 samples, three datasets were formed with 28

samples, and an additional dataset was formed with 50 samples. Each dataset was sampled from within the bounds set below in Table 8.

Table 8. Upper and lower boundaries for the sampled datasets used in the ESAV space mapping application

variable	bounds		units
	upper	lower	
area	30,000	25,000	in ²
	<i>208.33</i>	<i>173.61</i>	<i>ft²</i>
aspect ratio	4.0	2.0	-
t/c @ root	10	6	%
wing sweep	45	25	deg
taper ratio	0.5	0.2	-
t/c @ tip	8	4	%
camber location	10	0	%
max camber location	10	0	%
wing twist	5	-10	deg

Even after the development and debugging cases shown in Chapter IV, an additional modification to the algorithm was needed for the space mapping application to the ESAV tools to be successful. The responses for the two models using the same shared design variables differ by a magnitude of approximately 3,500 pounds within this region of the design space, and this difference in the magnitude of the responses is not correlated with any of the design variables. The reason a large difference exists between the two models is not relevant with regards to this space mapping technique (these reasons are known to the researchers in AFRL/RQ), but the space mapping algorithm needs to handle such offsets should they exist. Without modifying the algorithm, the only method for the space mapping process to handle such an uncorrelated offset is to increase the magnitude of the last coefficient in the space mapping form (for both the Task 3A and 3B forms). This last coefficient, or the “offset” coefficient, is the C_0 term in Equation 6. Table 9 shows the actual space mapping coefficients for a first-order least-squares fitting using the space mapping algorithm described in Chapter III. Notice the relative size of the offset coefficients compared to the other coefficient values.

The large offset coefficients are a result of how the minimization process in Task 2 handles the offset between the fidelity levels. To minimize the errors between the two models, all of the low-fidelity inputs needed significant scaling to make up for the inherent magnitude difference between the fidelity levels. Through the least-squares fitting of this data to the first-order form

Table 9. Space mapping coefficient values for the ESAV application without the modification to the algorithm

		scaled high-fidelity design variables									
		$\bar{x}_{H_1}$	$\bar{x}_{H_2}$	$\bar{x}_{H_3}$	$\bar{x}_{H_4}$	$\bar{x}_{H_5}$	$\bar{x}_{H_6}$	$\bar{x}_{H_7}$	$\bar{x}_{H_8}$	$\bar{x}_{H_9}$	offset
scaled low-fidelity design variables	$\bar{x}_{L_1}$	0.92	0.05	0.10	-0.22	0.02	-0.02	0.01	0.02	0.03	1.23
	$\bar{x}_{L_2}$	-0.53	-0.42	0.84	-1.21	0.15	0.11	-0.02	0.18	0.15	5.98
	$\bar{x}_{L_3}$	-0.01	-0.18	1.17	0.46	0.00	0.02	-0.01	-0.06	0.05	-2.17
	$\bar{x}_{L_4}$	0.15	0.29	0.20	1.15	-0.03	-0.08	0.05	-0.09	-0.01	2.59
	$\bar{x}_{L_5}$	-0.02	-0.03	0.04	-0.06	0.99	-0.00	0.00	0.01	0.01	0.27

shown in Table 9, the space mapping algorithm accomplished this scaling primarily through the offset coefficients. While this scaling was necessary to match the model responses, the design contours of the low-fidelity model were skewed as a result. This led to surrogate models that, while able to approximate the high-fidelity response, consistently under-performed in terms of accuracy with respect to more traditional PRM surrogates. In short, any potential gains in prediction accuracy due to similarities in the contours of the shared design space were lost by the space mapping algorithm accounting for the response difference through the scaling of design variables.

The modification made to the space mapping algorithm to alleviate this problem is simple, and yet it carries both an additional assumption and a procedural penalty with it. Before the minimization sequences in Task 2, the average difference between the high and low-fidelity responses using the same shared design variables is calculated. This average error is then added to the low-fidelity response when computing the objective function in Equation 35, which results in the following equation to replace the objective function in the Task 2 minimization sequences:

$$\min_{\bar{x}_L} J = [R_H(\bar{x}_H) - R_L(\bar{x}_L) + R_{\text{avg diff}}]^2. \quad (50)$$

Adding the average difference between the fidelity levels to the objective function assumes that this difference value is not attributable to any of the high-fidelity design variables. This assumption is deemed to be valid in this case due to the presence of large offset coefficient values in the space mapping matrix. Table 10 shows the space mapping coefficient values for the same data points used in constructing Table 9 after this modification to the algorithm was implemented.

The procedural penalty associated with modifying the objective function, as shown in Equation 50, is the new requirement that the steps in Task 2 can only begin once all of the

Table 10. Space mapping coefficient values for the ESAV application with the modification to the algorithm

		scaled high-fidelity design variables									
		$\bar{x}_{H_1}$	$\bar{x}_{H_2}$	$\bar{x}_{H_3}$	$\bar{x}_{H_4}$	$\bar{x}_{H_5}$	$\bar{x}_{H_6}$	$\bar{x}_{H_7}$	$\bar{x}_{H_8}$	$\bar{x}_{H_9}$	offset
scaled low-fidelity design variables	$\bar{x}_{L_1}$	0.96	-0.08	0.02	-0.09	0.03	0.02	0.00	0.07	-0.01	-0.00
	$\bar{x}_{L_2}$	-0.18	0.60	0.15	-0.46	0.10	0.13	0.01	0.28	-0.10	0.05
	$\bar{x}_{L_3}$	0.06	0.13	0.96	0.17	-0.04	-0.05	-0.01	-0.11	0.03	0.01
	$\bar{x}_{L_4}$	-0.07	-0.17	0.05	0.80	0.07	0.03	0.01	0.15	-0.02	-0.03
	$\bar{x}_{L_5}$	-0.01	-0.02	0.00	-0.02	1.01	0.00	-0.00	0.02	-0.00	0.00

high-fidelity responses have been gathered in Task 1. In the original space mapping algorithm, it was possible to execute the minimization sequence associated with each high-fidelity response as soon as the high-fidelity computation was complete. Now, with the need to calculate the average difference between the fidelity levels, the minimization sequences cannot begin until the last high-fidelity response is delivered. The inclusion of this average difference between fidelity levels alters the implementation of the surrogate model, as shown in Figure 37.

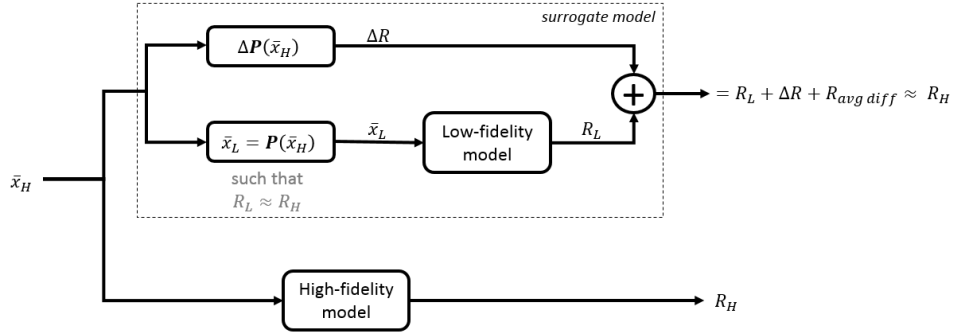


Figure 37. Final representation of the surrogate model constructed through the implementation of the space mapping algorithm

Due to the expense of running the high-fidelity analysis and the increased number of additional design variables, a qualitative analysis using surface responses (as shown in Chapter IV) is infeasible in this case. Instead, a stochastic analysis of the performance of the resulting surrogate models is performed. Each space mapping requires a certain number of sampled data points to construct a surrogate; the remaining number of sampled points can be used to determine the performance of the space-mapped surrogate in predicting the high-fidelity response. The number of remaining sample points in all cases is large enough to be considered statistically

relevant (greater than 40 members in the sample population). This large pool of data is used to conduct a stochastic analysis to characterize the accuracy of the surrogate models within the bounds specified in Table 8.

Consider the case where the first dataset (containing 14 data points) is used to derive a first-order least-squares surrogate model using the algorithm laid in Task 3A. The remaining data points are then used for comparison purposes to gather the percent errors between the high-fidelity response and the surrogate response. This data is analyzed through the use of histograms (to illustrate the number of occurrences within a set range of errors) as well as representative normal distribution curves generated using the mean and standard deviation of the percent errors collected. For comparison purposes, a traditional polynomial response surrogate is also generated using the same dataset as the space mapping implementation. This polynomial response surrogate was constructed using the same least-squares technique (detailed in Chapter II) that is used to fit data in the space mapping algorithm. The results for the traditional polynomial response surrogate are plotted alongside the results for the space-mapped surrogate, where applicable.

ESAV Space Mapping Results

The surrogate construction methods detailed in this thesis proved capable of producing surrogate models to approximate the high-fidelity response for each of the space mapping forms. In the sections that follow, the surrogates that resulted from each space mapping form are discussed and compared to the appropriate least-squares PRM surrogate models. The results from these comparisons with PRM surrogate models suggest that modest increases in estimation accuracy are possible through implementation of the space mapping algorithm detailed in this document. While the results presented in the following sections are not an exhaustive investigation of this alternative method for surrogate construction, they should be considered as a proof-of-concept for a method that might be beneficial within a subset of multifidelity design frameworks.

First-Order Datasets

Three of the seven datasets sampled in this research contain 14 data points and were dedicated to the construction of first-order surrogate models (with the exception that the nonlinear space mapping form discussed in Task 3B can be second-order). The number of sample points in each first-order dataset precludes the option of a kriging implementation. The results from the

least-squares space-mapped surrogate model using the first dataset show that this surrogate model is capable of approximating the high-fidelity response to a good degree. When the remaining sample points are fed into this surrogate model for comparison to the high-fidelity data, the surrogate model is able to predict the high-fidelity response to within a couple of percentage points of the total value. Figure 38 shows a histogram of the percent errors as well as a representative normal distribution plotted using the mean percent error and the standard deviation from the remaining high-fidelity data points.

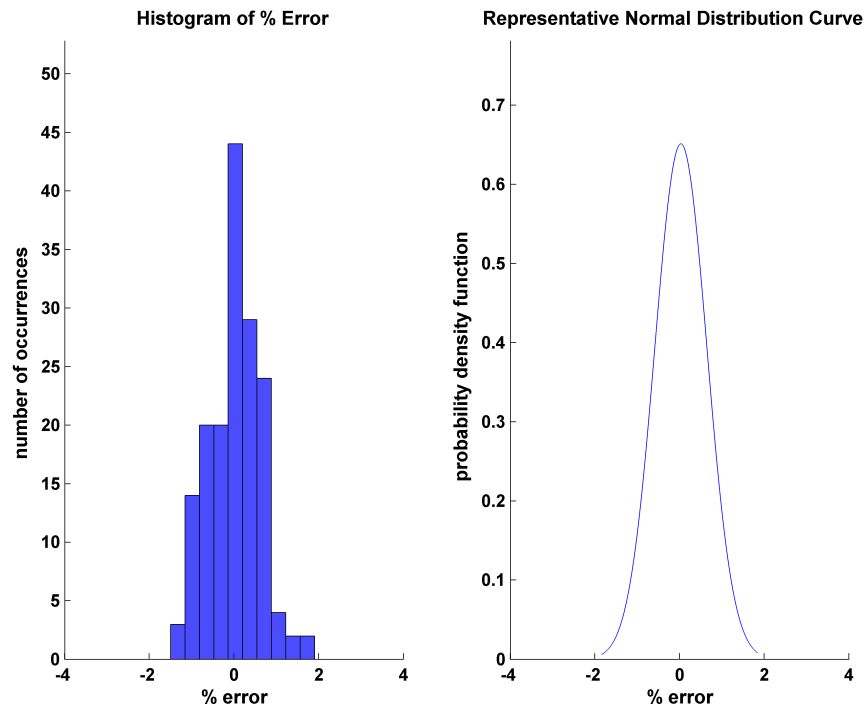


Figure 38. Histogram and representative normal distribution curve for the percent errors found using the least-squares space-mapped surrogate model in comparison to the high-fidelity response

The nonlinear space mapping form detailed in Task 3B produced a surrogate model with even greater accuracy in the prediction of results. The histogram and representative normal distribution curve are shown in Figure 39. The greater prediction accuracy is likely a result of the nonlinear nature of this space mapping form conforming more easily to the data obtained from the minimization sequences of Task 2 in the algorithm. Even though the number of sample points in this dataset are only sufficient for a linear least-squares fitting, the nonlinear form (found through the use of a genetic algorithm) is able to fit each design variable with either a linear or a second-order polynomial. While the resulting surrogate is more accurate than the linear

least-squares surrogate shown in Figure 38, the process of determining the space mapping coefficients takes significantly longer. This length of time depends upon the size of the initial population, the number of generations, and many other configurable options available in the genetic algorithm.

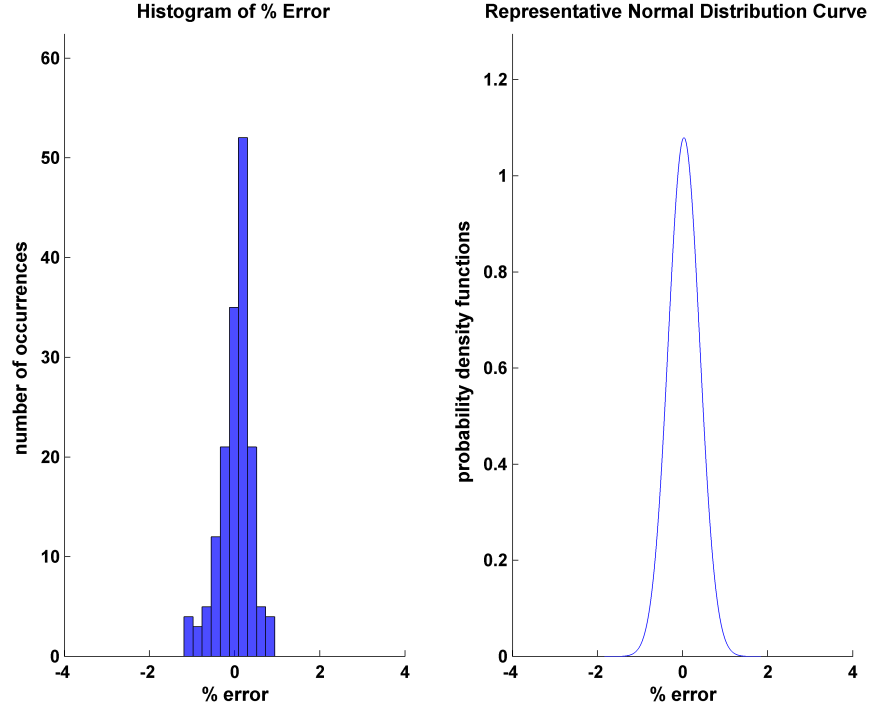


Figure 39. Histogram and representative normal distribution curve for the percent errors found using the nonlinear space-mapped surrogate model in comparison to the high-fidelity response

Using the same sample datapoints in the dataset, a first-order least-squares PRM surrogate was constructed and the remaining high-fidelity design vectors were fed into this surrogate model for a similar error analysis. Both the least-squares and nonlinear space-mapped surrogate models display reduced spread in the range of the percent errors with respect to the PRM surrogate, and the representative normal distribution curves show a similar advantage for the space-mapped surrogates in terms of prediction accuracy as well. These comparisons are shown in Figure 40.

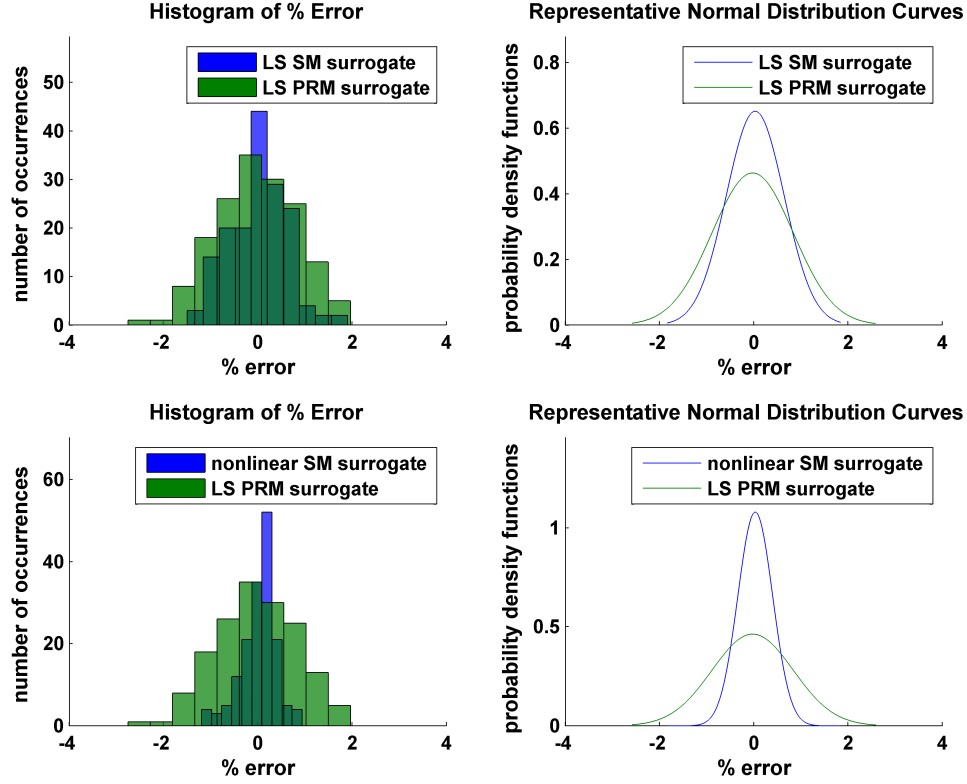


Figure 40. Histogram and representative normal distribution curve for the percent errors found using the polynomial response surrogate (LS PRM) overlaid on the data from the space-mapped (SM) surrogate

A scatterplot of the actual responses for both the high-fidelity model and the surrogate models is shown in Figure 41. The dotted black line in each of the plots signifies a perfect relationship between the high-fidelity responses and the surrogate response. Data points off of this line therefore have error associated with the surrogate's prediction. While a scatterplot provides a good visualization of the data, a quantitative determination of which surrogate is more accurate can be obtained through a RMSE calculation for each of the data points shown. The RMSE comparisons for each of the first-order datasets can be found in Table 11.

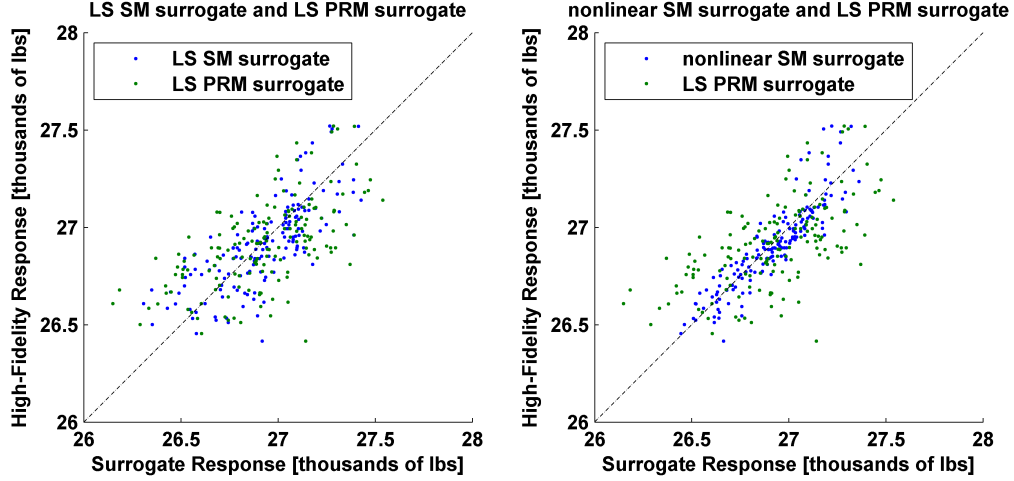


Figure 41. Scatterplot for both the least-squares and nonlinear space-mapped surrogate responses, with the least-squares PRM surrogate response plotted for comparison

Applying the same process that resulted in the data shown in Figure 40 for the remaining two datasets shows a more equal footing between the space-mapped surrogate models and the least-squares PRM surrogate constructed from the same dataset. Considering the information shown in Figures 40, 42, and 43, the space-mapped surrogate models are at least as capable as their PRM surrogate counterparts in estimating the high-fidelity response.

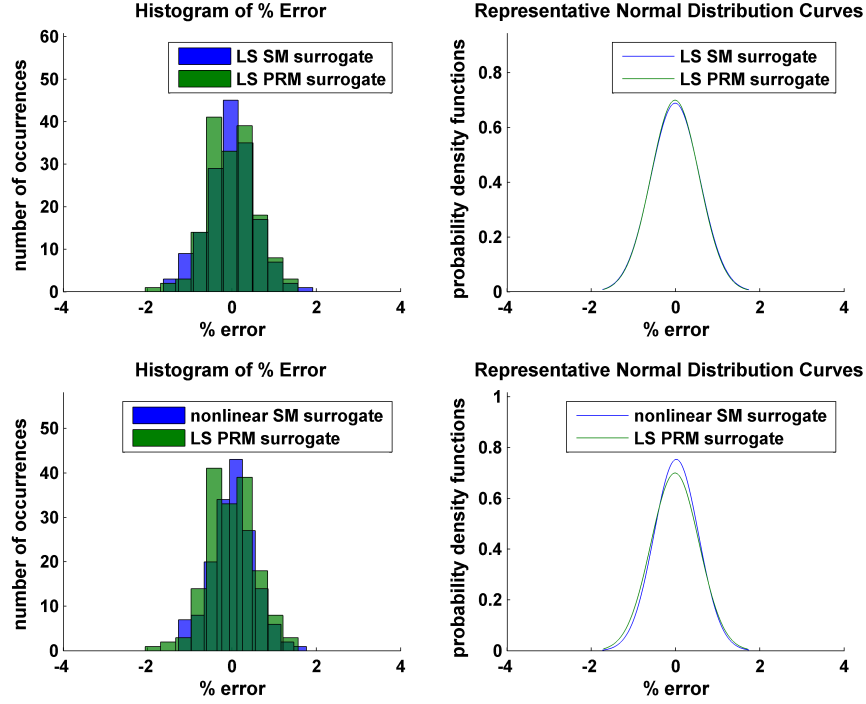


Figure 42. Percent error comparison between the least-squares space-mapping and the PRM surrogate models derived from samples in the second dataset

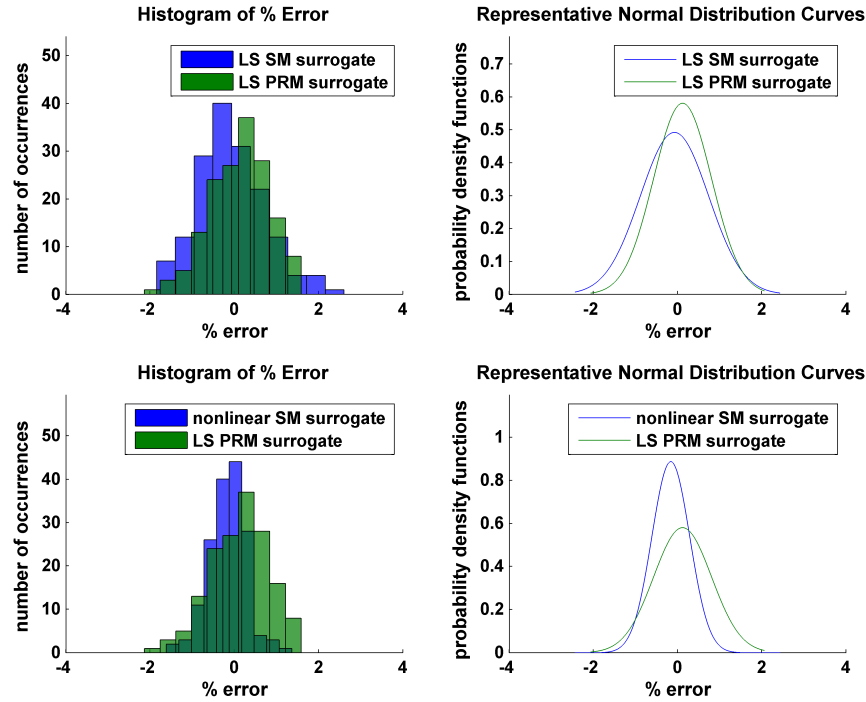


Figure 43. Percent error comparison between the least-squares space-mapping and the PRM surrogate models derived from samples in the third dataset

The following table contains a summary of the comparison results for the first-order datasets. The RMSE shown in Table 11 was calculated using the percent errors (as opposed to the actual error numbers). The RMSE is included as an additional comparison parameter and in all three cases the lowest RMSE corresponds to the surrogate model with the least standard deviation. Scatterplots for the second and third first-order dataset have been included in the Appendix in Figures 58 and 59.

Table 11. Surrogate performance summary comparing the various surrogates constructed from the sampled datasets containing 14 data points

	surrogate type	average % error	standard deviation	max % error	min % error	RMSE*
dataset 1	1 st -order LS** SM [†]	0.0286	0.6129	1.9020	-1.4972	0.6117
	1 st -order LS PRM [‡]	-0.0261	0.8615	1.9652	-2.7509	0.8593
	nonlinear polynomial SM	0.0294	0.3698	0.9413	-1.1912	0.3698
dataset 2	1 st -order LS SM	-0.0122	0.5797	1.9157	-1.6243	0.5780
	1 st -order LS PRM	-0.0160	0.5706	1.5699	-2.0582	0.5691
	nonlinear polynomial SM	0.0134	0.5299	1.7668	-1.2683	0.5285
dataset 3	1 st -order LS SM	-0.0751	0.8116	2.6026	-1.8438	0.8125
	1 st -order LS PRM	0.1170	0.6878	1.5849	-2.1322	0.6956
	nonlinear polynomial SM	-0.1621	0.4503	1.3652	-1.6147	0.4773

* root mean square error using the % error, ** least-squares

[†]space mapping, [‡]polynomial response methodology

Second-Order Datasets

In order to fit a second-order least-squares form to the space mapping data, three datasets were sampled with 28 sample points in each. Surrogate models were constructed using these sample points in the same manner as before which enabled the comparison of the space-mapped surrogate models with the least-squares PRM surrogates. The surrogates (both the space-mapped surrogates and the PRM surrogates) based on the second-order polynomial forms do not necessarily improve the predictive accuracy of the model in comparison to the surrogates based on first-order polynomials. This finding implies that the high-fidelity response, within the bounds of the design space set in Table 8, is predominantly linear with respect to the design variables. For

each of the three second-order datasets, the space-mapped surrogate models perform at least as well as the least-squares PRM surrogates constructed from the same data points.

The surrogates built from the data points in the first dataset were found to predict the high-fidelity response to within a small range of percent errors. A comparison between the surrogate models based upon a least-squares fitting of data is shown in Figure 44. Judging from the information available in the graphic, there does not appear to be a clearly-superior surrogate model with respect to predictive accuracy. A similar comparison between the second-order least-squares PRM surrogate and the nonlinear space-mapped surrogate is shown in Figure 45. The nonlinear space mapping form yielded a surrogate with similar predictive accuracy as the second-order least-squares PRM surrogate. Additional scatterplots comparing each of the surrogate methods for each of the three datasets are included in the Appendix in Figures 64-69.

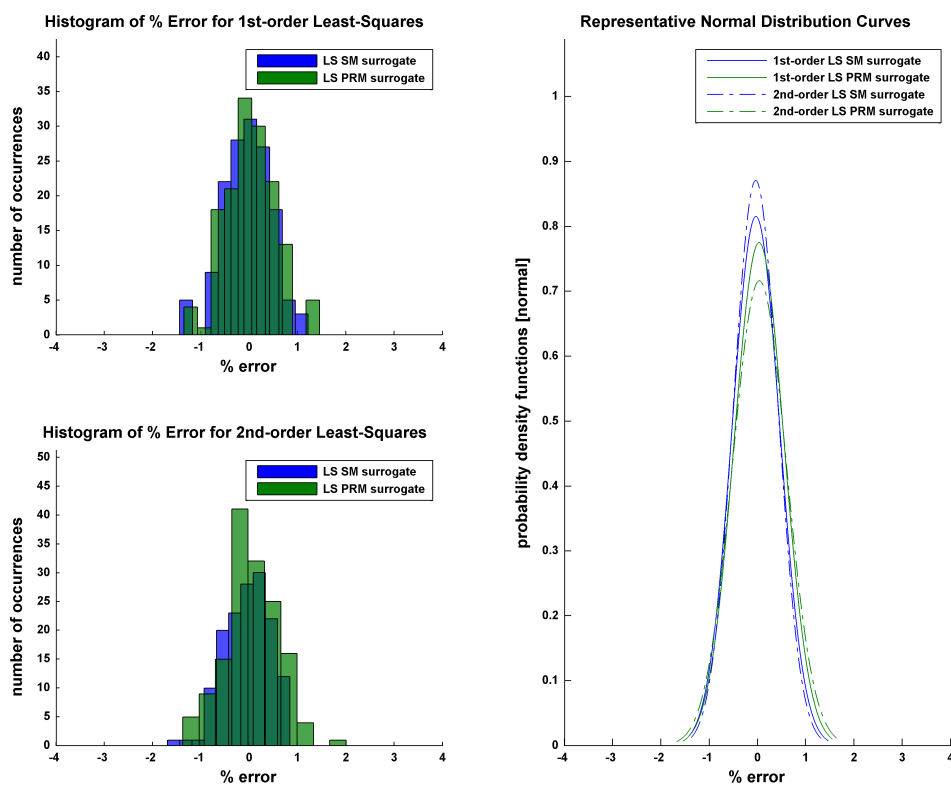


Figure 44. Histograms and representative normal distribution curves for the first and second-order least-squares surrogate models based on the first of three second-order datasets

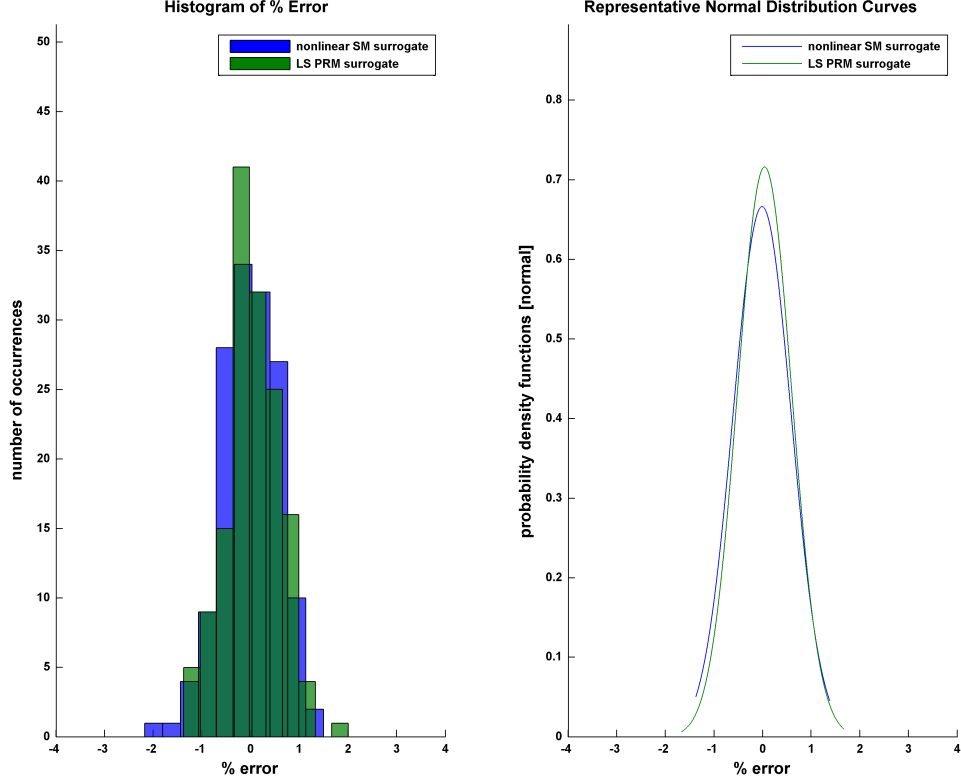


Figure 45. Histogram and representative normal distribution curve for the nonlinear surrogate model based on the first of three second-order datasets (second-order LS PRM surrogate data plotted for comparison)

Illustrations similar to Figures 44 and 45 were produced for the second and third dataset and can be found in the Appendix in Figures 60-63. While there are differences between the predictive capacity of each surrogate model type across the three datasets, in general the space-mapped surrogate models proved capable of approximating the high-fidelity response equally as well as the least-squares PRM surrogate models. In many of the cases, the space-mapped surrogates yielded modest gains in predictive accuracy over their least-squares counterparts. Table 12 contains the important data obtained from the stochastic analysis of each of the surrogate models discussed. Included for each surrogate model is the RMSE value, which is yet another means of comparison between the many surrogate models. In all three of the second-order datasets, a space-mapped surrogate yielded the smallest RMSE.

Table 12. Surrogate performance summary comparing the various surrogates constructed from the sampled datasets containing 28 data points

	surrogate type	average % error	standard deviation	max % error	min % error	RMSE*
dataset 1	1 st -order LS** SM [†]	-0.0352	0.4896	1.2166	-1.4397	0.4892
	1 st -order LS PRM [‡]	0.0299	0.5150	1.4534	-1.3507	0.5141
	2 nd -order LS SM	-0.0396	0.4584	0.8378	-1.6909	0.4585
	2 nd -order LS PRM	0.0369	0.5573	2.0011	-1.3740	0.5567
	nonlinear polynomial SM	-0.0143	0.5990	1.4974	-2.1739	0.5972
dataset 2	1 st -order LS SM	0.0371	0.3835	1.1081	-1.0523	0.3840
	1 st -order LS PRM	-0.0493	0.4174	1.1670	-1.1908	0.4189
	2 nd -order LS SM	0.0477	0.5339	1.5520	-1.7420	0.5343
	2 nd -order LS PRM	-0.0562	0.5212	1.6288	-1.5705	0.5225
	nonlinear polynomial SM	0.1039	0.4837	1.5385	-1.1228	0.4931
dataset 3	1 st -order LS SM	0.1337	0.4762	1.8175	-0.8584	0.4931
	1 st -order LS PRM	-0.1438	0.5486	0.9258	-2.0655	0.5653
	2 nd -order LS SM	0.1294	0.5350	1.4200	-1.0987	0.5486
	2 nd -order LS PRM	-0.1446	0.6487	1.4751	-1.6820	0.6625
	nonlinear polynomial SM	0.0973	0.5544	1.8773	-1.0603	0.5610

* root mean square error using the % error, ** least-squares

[†]space mapping, [‡]polynomial response methodology

Nonlinear Kriging Space-Mapped Surrogates

The construction of surrogate models based on Kriging interpolation schemes requires a dataset of at least 58 sample points. Since none of the datasets referenced up to this point have the requisite number of points, multiple datasets were merged so that two instances of kriging-based surrogate comparisons could be made. The first kriging dataset is formed from the dataset containing 50 sample points added to the first dataset of 14 sample points, bringing the number of sample points to 64. The second kriging dataset was constructed from the first and

second datasets containing 28 samples as well as the second dataset containing 14 samples points (for a total of 70 sample points).

The performance of the kriging-based space-mapped surrogate models was found to be similar to the performance of the traditional kriging surrogate models constructed from the same sampled points. A comparison between the two surrogate models using the bare minimum of 58 data points was conducted and the results are shown in Figure 46. This comparison put to the test the hypothesis that the similarity between the contours of the shared design space might significantly improve the predictive capability of the space-mapped surrogate model for a sparse number of points. The results do not confirm this theory, as shown in the similarity between the two histograms and representative normal distributions.

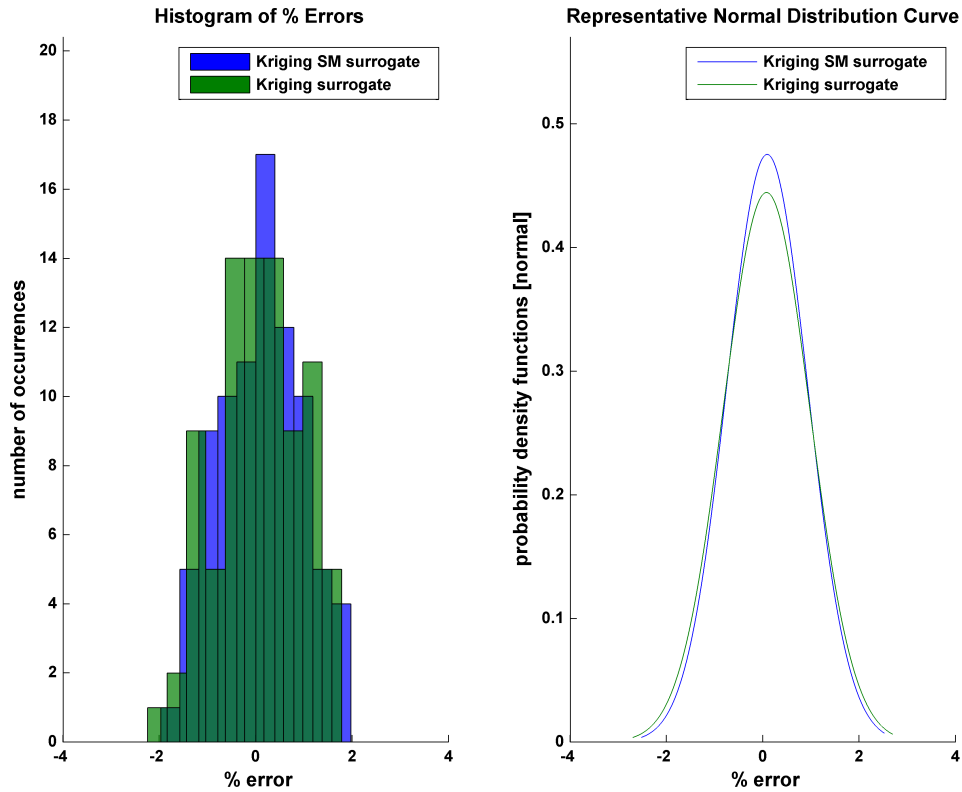


Figure 46. Histograms and representative normal distribution curves for the kriging-based SM surrogate and the traditional kriging surrogate based on the same 58 sample data points from the first kriging dataset

Plotting the RMSE against the number of sample points dedicated to the surrogate construction shows the the predictive capabilities for each surrogate. While the kriging

space-mapped surrogate displays slightly smaller RMSE values in comparison to the traditional kriging implementation, the difference between the surrogate models is deemed to be negligible. The kriging space-mapped surrogate model, for this model pairing, did not produce the significant accuracy gains to justify the added complexity of the space mapping algorithm versus a traditional kriging method.

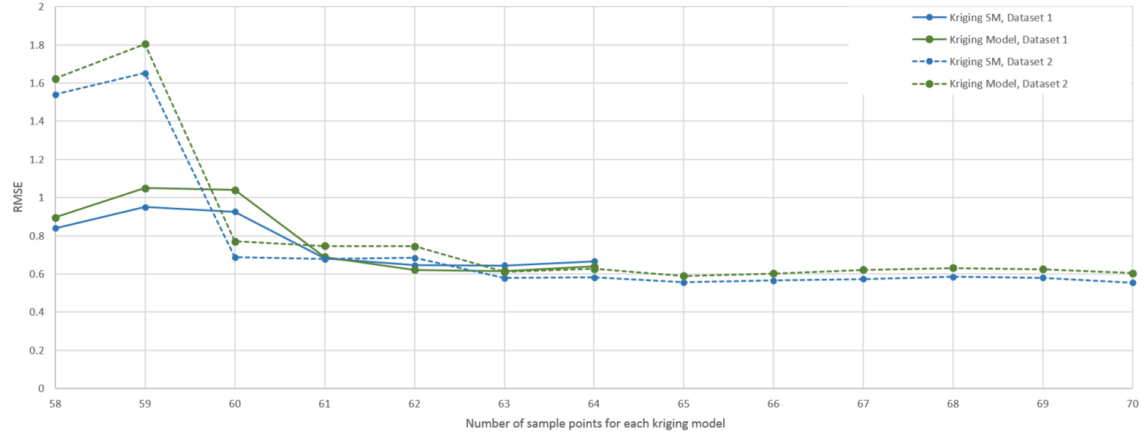


Figure 47. RMSE plotted against the number of samples on which each kriging surrogate was based

VI. Conclusions

The intent of this research was to explore the possibility of realizing gains in predictive accuracy of a surrogate model over existing surrogate methods through the alignment of a low-fidelity design space with that of the high-fidelity design space using space mapping techniques. Accuracy in a surrogate model is important because the ability to accurately approximate the high-fidelity design space contours means a design team could access high-fidelity information about the system at the computational expense of the low-fidelity models. This modeling capability would serve to mitigate the risk identified by AFRL/RQ that the lower fidelity design tools might exclude novel design configurations able to leverage new technologies or relevant physical phenomena.

The methods employed in the space mapping algorithm yielded surrogate models that, in the majority of the cases explored in Chapter V, met or exceeded the predictive capability of the traditional methods of surrogate construction. In the cases where a space-mapped surrogate model was the more accurate surrogate in the prediction of the high-fidelity response, the gains in accuracy may not be large enough to warrant the added complexity associated with the space mapping algorithm. The process presented in this work is a first attempt to build a surrogate model using the modified space mapping procedure outlined in Chapter III, and it is possible that future revisions and modifications might take better advantage of the information available at the lower fidelity levels to improve the accuracy of the surrogate models.

There are both advantages and disadvantages to the employment of space mapping in the construction of surrogate models. Although the results shown in this document are only a proof of concept and the method has yet to be tested in a large number of applications to discover how robust the technique truly is, the technique did show modest gains in predictive accuracy for the ESAV application. It is therefore possible that the employment of the space mapping technique can result in a more accurate surrogate model. Any gains in predictive accuracy will come at a price, however. Any surrogate construction technique will require a certain number of points to be sampled from the high-fidelity model; this space mapping algorithm also calls for multiple optimization processes which require the execution of the low-fidelity model as well. Additionally, each prediction on the part of the space-mapped surrogate will require the execution the low-fidelity model. Depending upon the application, the costs of executing the low-fidelity tool so often may outweigh any increase in the predictive capability of the surrogate model. The

traditional PRM surrogate techniques, on the other hand, only require the high-fidelity data points. Execution of the surrogate model in this case is simply the evaluation of the polynomial.

The development of this alternative method for approximation proceeded in the hopes of finding a viable method of using the information available at the lower fidelity levels to reduce the amount of information needed from the higher fidelity levels in the formation of a surrogate model. For the process to be useful, the resulting surrogate models need to be more accurate than conventional options built upon the same underlying data. This would bolster the argument that a space-mapped surrogate can achieve similar accuracy levels with less sample points. Results from this research do not show this method to be a significant improvement over current techniques used in the design community, but these initial findings justify further research into the incorporation of space mapping techniques in the field of surrogate construction.

Future Work

The following sections outline several avenues for future research that might better characterize the benefits space mapping techniques can bring to the field of surrogate construction. In addition to these research opportunities, a list of changes to the algorithm are presented that were conceived in the course of the research, but not explored. In general, this research would benefit greatly from a much wider application of the space mapping algorithm to many different model-pairings. This would allow for a better characterization of any gains in predictive accuracy resulting from the space mapping techniques.

Space Mapping as a Means to Reduce the Sampling of the High-fidelity Model

Suppose the premise of this research is true, and surrogate models constructed through space mapping techniques do provide more accurate approximations of the high-fidelity response than do contemporary surrogates. It would therefore be possible to achieve the same levels of predictive accuracy shown in contemporary surrogates by using a surrogate based on less high-fidelity sample points coupled with a low-fidelity analysis through space mapping techniques. This process would lessen the number of samples required from the high-fidelity model, which would lessen the computational burden associated with the construction of the surrogate. A future task to test this hypothesis would need to implement a space mapping algorithm to a wider set of model-pairings spanning many different engineering disciplines. Many different datasets would need to be gathered

and a range of surrogate models (space-mapped and traditional) would then be constructed, evaluated against other high-fidelity data points, and compared over various values of q .

Characterizing the Relationship between Fidelity Levels

One characteristic of the space mapping process not actively researched in this work was an analysis of the information contained in the resulting space mapping forms. Information about the relationship between the high and low-fidelity models is present in the space mapping relationships and could be explored further. For instance, with regards to the ESAV application in Chapter V, the first-order least-squares space mapping matrix in Table 10 shows how strongly the impact of the shared design variables on the low-fidelity response correlate with the impact of the high-fidelity design variables on the high-fidelity response. If the two computational models share similar design space contours, then the shared design variables present in the low-fidelity model should correlate most strongly with their counterparts in the high-fidelity model. This is seen in Figure 48 for the majority of the shared design variables, with the exception of the aspect ratio (AR) of the wing.

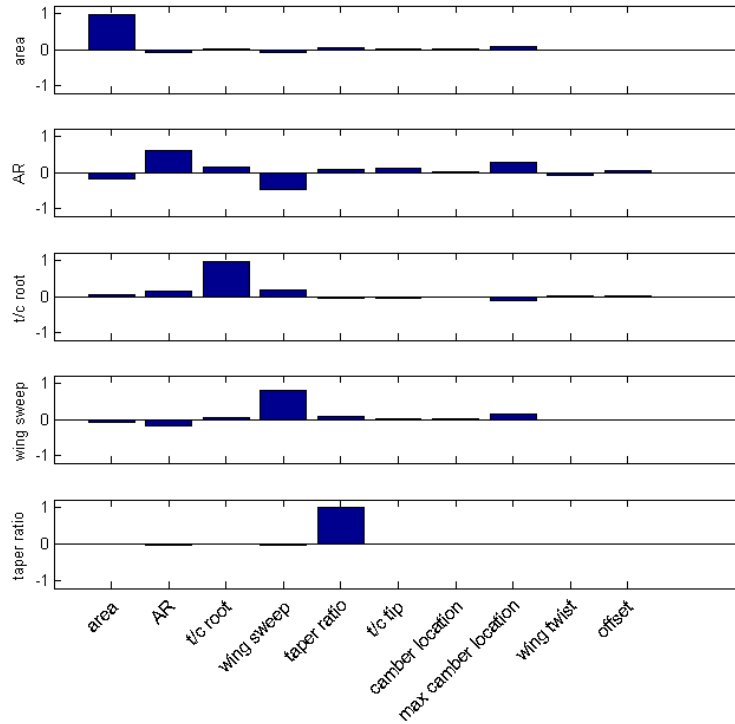


Figure 48. Bar chart illustrating the coefficient values of the space mapping form shown in Table 10

The data in Figure 48 suggests that the trends seen in the low-fidelity model will match the trends seen at the higher fidelity level, except in the case of the aspect ratio variable. The presence of relatively large coefficient values for design variables other than the high-fidelity aspect ratio variable in the space mapping of the low-fidelity aspect ratio variable indicate that changes in the aspect ratio value do not affect the responses of each model in the same way. Put another way, the influence of the additional design variables is most easily seen in the space mapping of the low-fidelity aspect ratio.

Potential Improvements to the Algorithm

In the course of this work, several minor alterations to the implementation of the space mapping algorithm were considered but not explored due to time constraints. The first proposed change is to the method of gathering the residual differences, ΔR . In the algorithm's current state, this quantity is calculated once the optimizer has changed the low-fidelity design variables in its attempts to match the high and low-fidelity responses. This process does not likely fulfill the intent of capturing the localized differences between fidelity levels, since the variable values in $\bar{x}_L^*$ could be significantly different than the shared design variable values, $\tilde{x}_L$. To better capture the localized difference, the algorithm should capture the residual as

$$\Delta R = R_H(\bar{x}_H) - R_L(\tilde{x}_L) - R_{\text{avg diff}} \quad (51)$$

so that the difference between fidelity levels is logged at the appropriate position in the shared design space, $\tilde{x}_L$. The optimal low-fidelity design vector, $\bar{x}_L^*$, output in this case should be $\tilde{x}_L$ so that the space mapping data for this specific high-fidelity design vector reflects the presence of a localized offset at the appropriate position in the shared design space.

The optimization process within Task 2, as executed in this research, was implemented without considering the costs associated with iteratively executing the low-fidelity analysis. While the low-fidelity model is assumed to be less expensive to execute than the high-fidelity model, this does not mean that the cost of using the low-fidelity model in an optimization scheme is going to be negligible. A refinement to the algorithm presented here would be to formulate the optimization process to reduce the number of low-fidelity analysis calls required to obtain the space mapping data.

The inclusion of the average offset between fidelity levels, $R_{\text{avg diff}}$, may alleviate the need for the offset coefficient in the space mapping form. The purpose of $R_{\text{avg diff}}$ is to remove the difference between the fidelity levels that is not attributable to the high-fidelity design variables; once removed, it may make some sense to attribute all other differences between the two models to the design variables. Removing the offset coefficient from the space mapping form would force the algorithm to attribute any remaining difference between the two fidelity levels to the design variables of the high-fidelity model.

Appendix

Nomenclature

<i>Traditional Space Mapping</i>	
symbol	definition
x	design variable
R	response of function or model, a function of design variable(s)
$\bar{x}_L$	low-fidelity design vector
n	number of low-fidelity design variables
$\bar{x}_H$	high-fidelity design vector
p	number of high-fidelity design variables
$\mathbf{P}$	space mapping relationship
R_L	low-fidelity response, function of $\bar{x}_L$
R_H	high-fidelity response, function of $\bar{x}_H$

<i>Least-Squares Projections and Polynomial Response Methodology</i>	
symbol	definition
$\mathbf{A}$	matrix consisting of high-fidelity design variables raised to the appropriate powers (see Equation 9)
$\mathbf{P}_\mathbf{A}$	projection matrix formed using $\mathbf{A}$ through Equation 3
k	degree of polynomial
e_{act}	actual error for a polynomial approximation, see Equation 11
e_{est}	estimated error for a polynomial approximation at a given location, see Equation 12
σ	standard deviation for the estimation error at a given location

Kriging

symbol	definition
$z^*(\bar{x})$	estimate at a location $\bar{x}$ in an interpolation scheme
q	number of sample data points
λ	weighting coefficient used in an interpolation scheme
$z(\bar{x})$	value of sample data point at the sample data location
$m(\bar{x})$	trend component of the sample data points
$\mathbf{K}$	covariance matrix used in the kriging process, see Equation 16
$r(\bar{x})$	residual component of the sample data points
$C_R(\bar{h})$	covariance between two sample points as a function of lag (h)
m	constant mean of the sample data points, used in simple kriging
$\bar{k}$	vector of covariance values between the sample data points and the estimation point

Modified Space Mapping

symbol	definition
$\tilde{x}_L$	design variables shared between fidelity levels
$\tilde{x}_H$	additional design variables in the high-fidelity model
ΔR	residual difference between fidelity levels at the end of the optimization sequence in Equation 35
$\bar{x}_L^*$	low-fidelity design vector that results from the optimization sequence in Equation 35
$\Delta \mathbf{P}$	approximation for the residual differences
μ_x	mean value of a sample population
σ_x	standard deviation of a sample population
E_t	total error calculation used in the evaluation of different space mapping forms
$\bar{\phi}$	vector of coefficients used in the nonlinear space mapping of Task 3B
$\bar{\beta}$	vector of powers used in the nonlinear space mapping of Task 3B
γ	scalar value used in the nonlinear space mapping of Task 3B

Variables in Model-Pairings

symbol	definition
ρ	density, units are slugs / ft ³
V	velocity, units are ft / sec
S	planform area, units are ft ²
$C_{L\alpha}$	lift-curve slope, units are 1 / deg
α	angle of attack, units are in degrees
eff	efficiency factor, unitless
AR	aspect ratio, unitless

Table 13. Variable values held constant in the production of Figure 12

density	velocity	planform area	lift-curve slope	efficiency factor	aspect ratio
0.002377	100	10	0.11	0.85	3.0
slug/ ft ³	ft/s	ft ²	deg ⁻¹	-	-

Table 14. Case 1 shared variable values for surrogate model comparisons

density	velocity	planform area	lift-curve slope	angle of attack
0.0023385	100	10	0.11	8
slugs/ft ³	ft/sec	ft ²	1/deg	deg

Table 15. Case 1 variable boundaries

<i>variable</i>	density	velocity	planform area	lift-curve slope	angle of attack	efficiency factor	aspect ratio
<i>upper</i>	0.002377	105	12	0.12	10	0.95	3.5
<i>lower</i>	0.002300	95	8	0.1	6	0.75	2
<i>units</i>	slugs/ft ³	ft/sec	ft ²	n/a	deg	n/a	n/a

Table 16. Case 1 variable boundaries

<i>variable</i>	density	velocity	planform area	lift-curve slope	angle of attack	$\bar{x}_{H_6}$	$\bar{x}_{H_6}$
<i>upper</i>	0.002377	105	12	0.12	10	2.356194	1.5
<i>lower</i>	0.002300	95	8	0.1	6	0.785398	0.5
<i>units</i>	slugs/ft ³	ft/sec	ft ²	n/a	deg	radians	n/a

Additional Data Plots and Illustrations

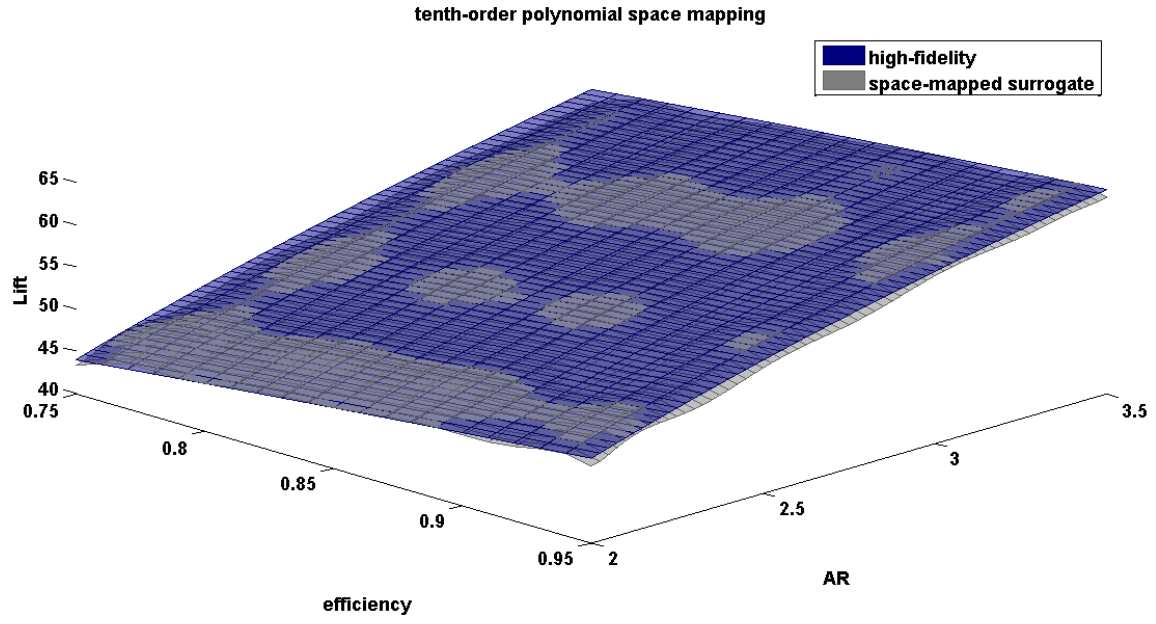


Figure 49. Tenth-order polynomial space mapped surrogate for the Case 1 model pairing

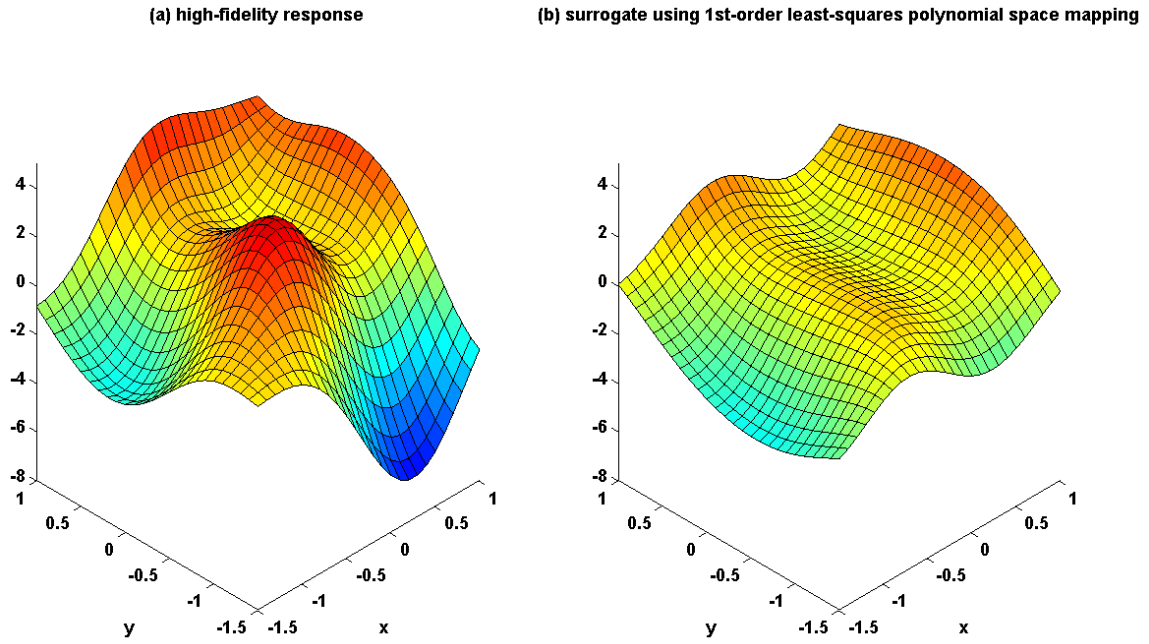


Figure 50. Surrogate constructed using a linear least-squares space mapping for the 3rd model pairing

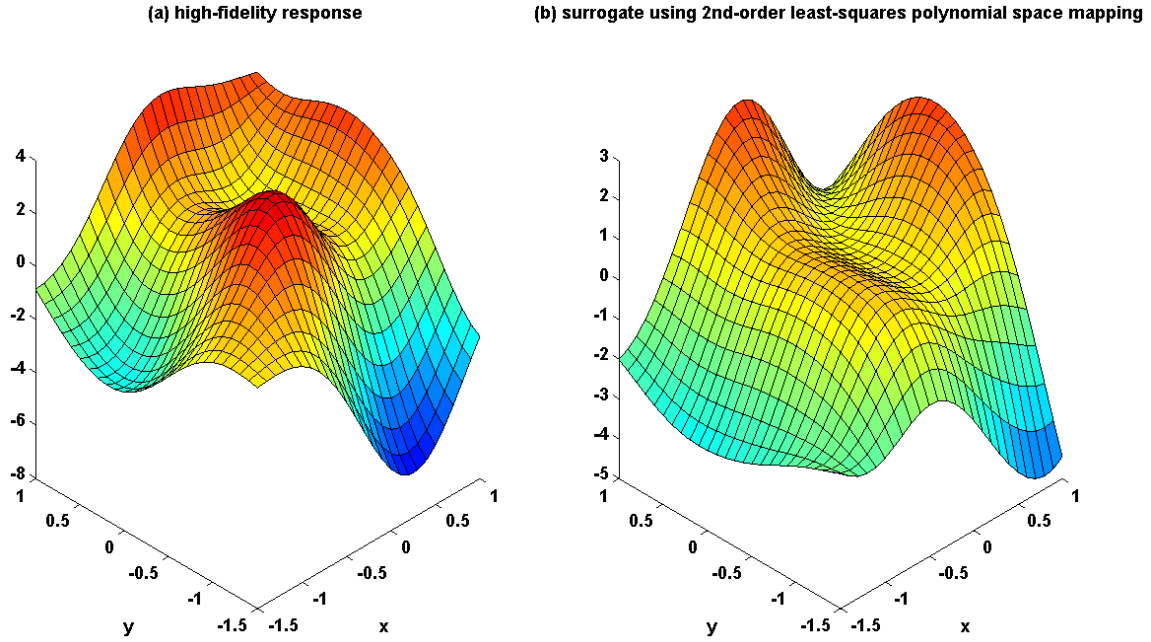


Figure 51. Surrogate constructed using a second-order least-squares space mapping for the 3rd model pairing

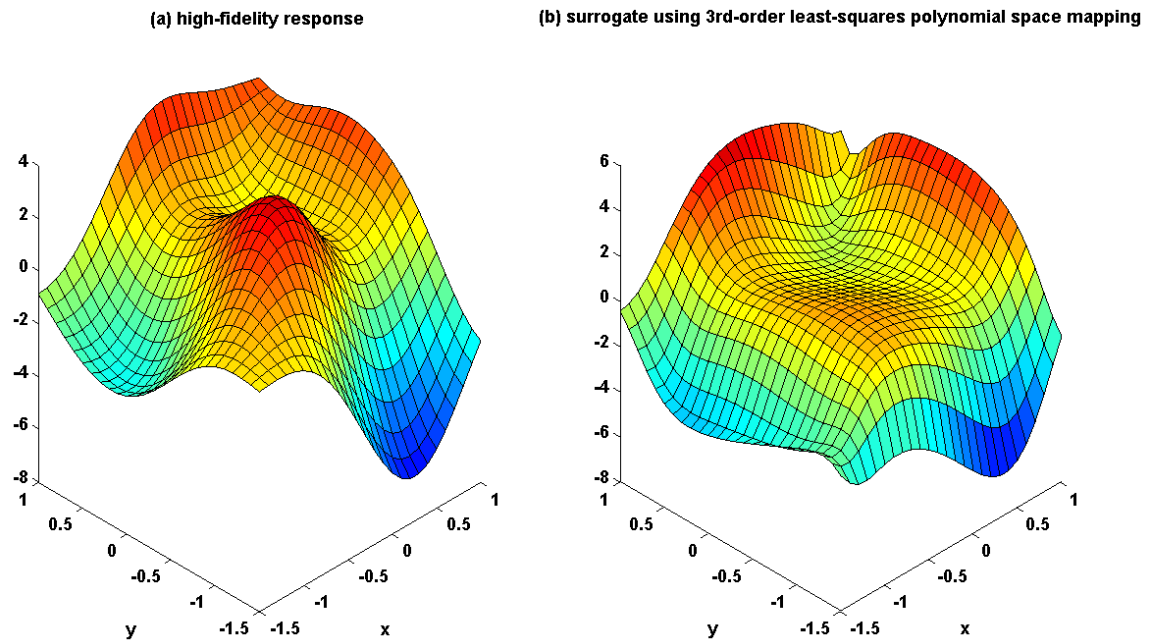


Figure 52. Surrogate constructed using a third-order least-squares space mapping for the 3rd model pairing

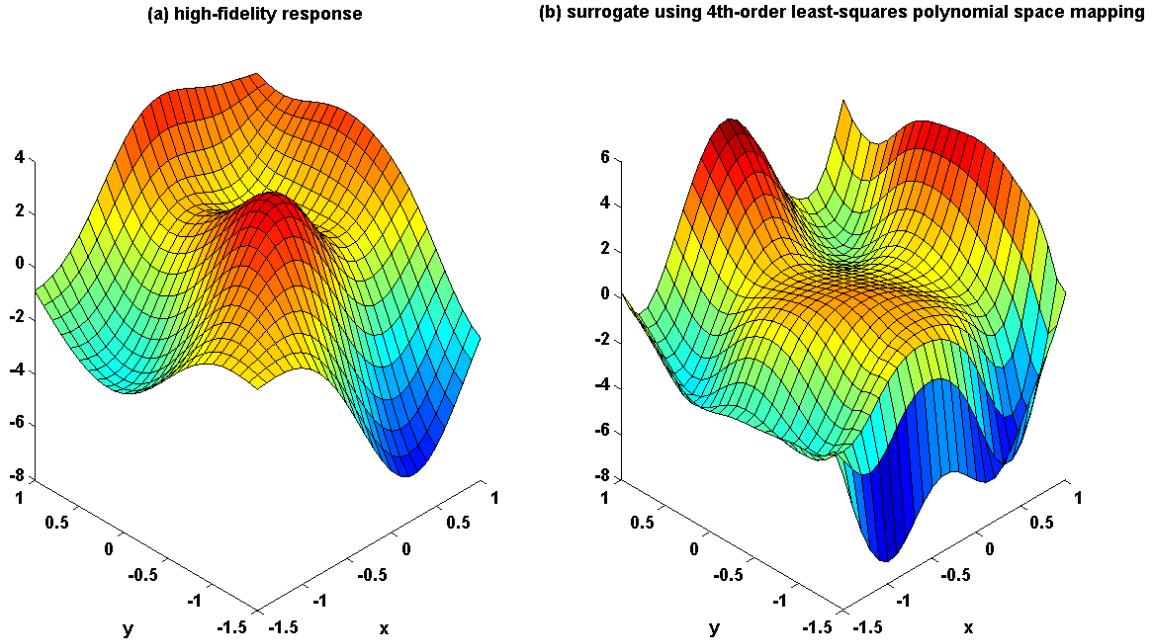


Figure 53. Surrogate model constructed using a fourth-order least-squares space mapping approach

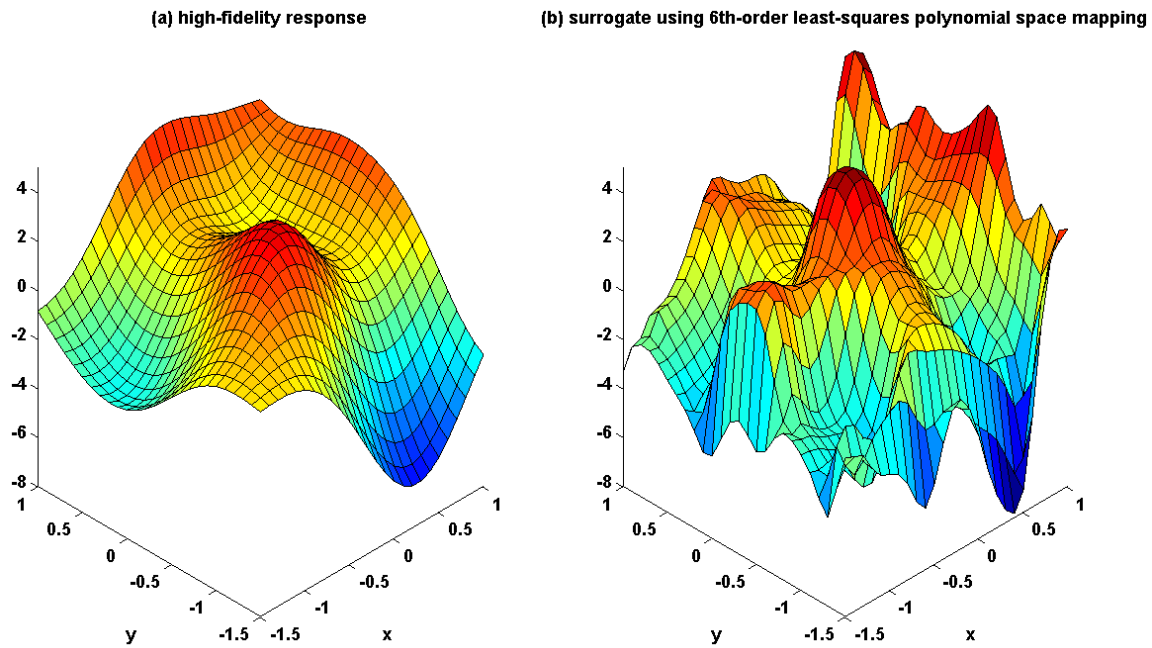


Figure 54. Surrogate model constructed using a sixth-order least-squares space mapping approach

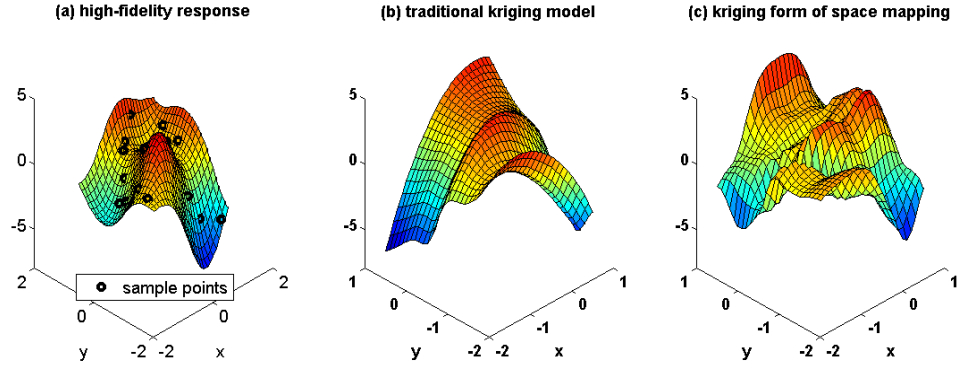


Figure 55. Comparison between a kriging-based space-mapped surrogate and the corresponding traditional kriging surface built from the sampling locations shown. RMSE values: (b) 2818.6 (c) 9679.2

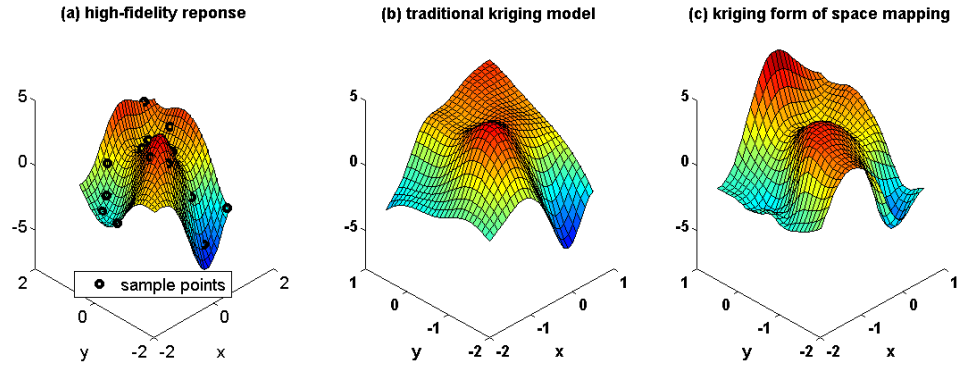


Figure 56. Comparison between a kriging-based space-mapped surrogate and the corresponding traditional kriging surface built from the sampling locations shown. RMSE values: (b) 645.1 (c) 910.6

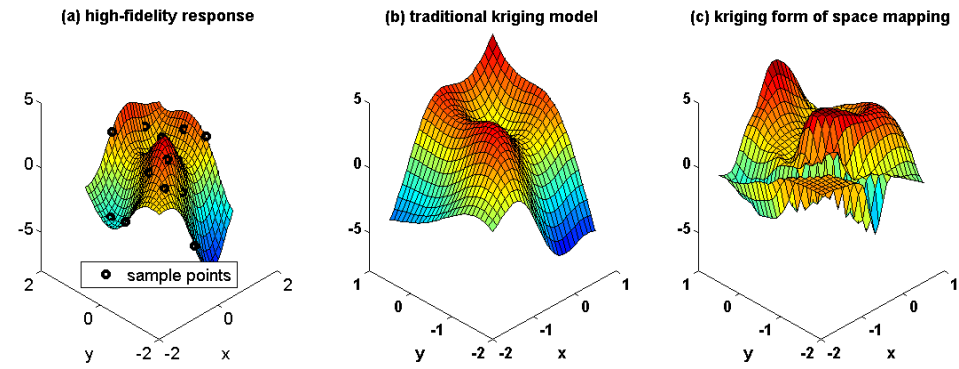


Figure 57. Comparison between a kriging-based space-mapped surrogate and the corresponding traditional kriging surface built from the sampling locations shown. RMSE values: (b) 459.0 (c) 2379.6

Table 17. Inputs held constant in the weight predictor for the ESAV space mapping

<i>input</i>	<i>value</i>	<i>units</i>
aspect ratio of vertical tail	0.9	-
number of vertical tails	1	-
horizontal tail span	18	ft
fuselage structural depth	6.5	ft
engine diameter	3.875	ft
fuselage width at horizontal tail intersection	7.5	ft
duct constant	3.43	-
fuselage structural length	49.5	ft
electrical routing distance (generators to avionics to cockpit)	40	ft
duct length	15	ft
length from engine front to cockpit	15	ft
single duct length	0	ft
length of engine shroud	16.5	ft
tail length	16.67	ft
length of tail pipe	1	ft
Mach number	2	-
crew number	1	-
number of crew equivalents	1	-
number of engines	1	-
number of generators	1	-
ultimate landing load factor	9	-
number of nose wheels	1	-
number of people on board	1	-
number of flight control systems	3	-
number of fuel tanks	7	-
number of hydraulic utility functions	10	-
ultimate load factor	11	-
system electrical rating	160	kV A
total area of control surfaces	200	ft ²
control surface area (wing-mounted)	75	ft ²
firewall surface area	2	ft ²
horizontal tail area	98	ft ²
rudder area	21	ft ²
vertical tail area	86	ft ²
specific fuel consumption	1.5	1/hr
total engine thrust	23,830	lbs
single engine thrust	23,830	lbs
integral tanks volume	900	gal
self-sealing (protected) tanks volume	17	gal
total fuel volume	1,076	gal
total fuselage structural width	7.5	ft
flight design gross weight	22,500	lbs
engine weight	3,067	lbs
landing design gross weight	19,500	lbs
uninstalled avionics weight	1,000	lbs
vertical tail sweep at 1/4 chord	45	deg
taper ratio for vertical tail	0.012	-

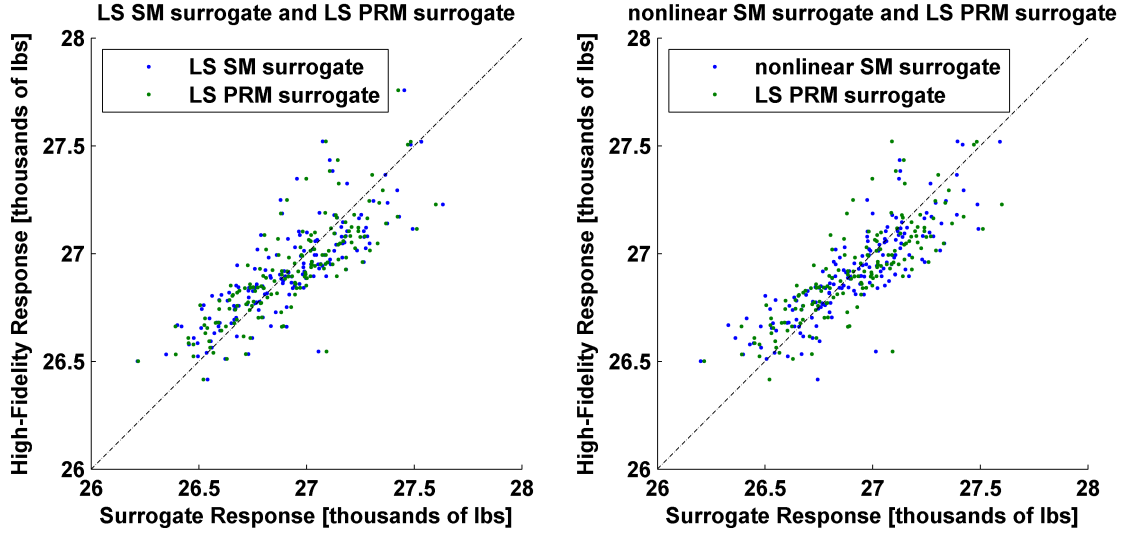


Figure 58. Scatterplot for both the least-squares and nonlinear space-mapped surrogate responses, with the least-squares PRM surrogate response plotted for comparison (second of three first-order datasets)

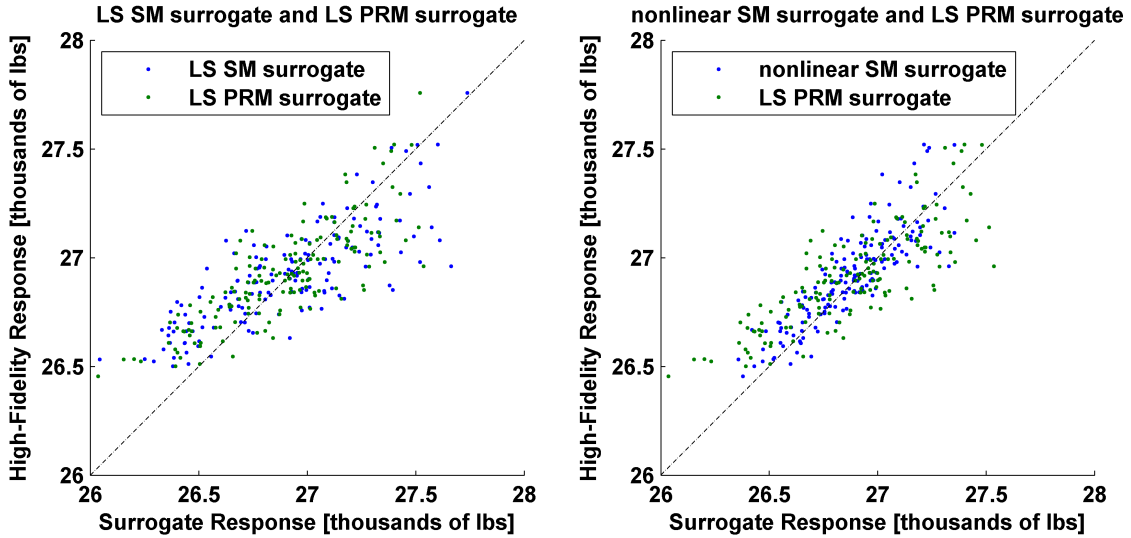


Figure 59. Scatterplot for both the least-squares and nonlinear space-mapped surrogate responses, with the least-squares PRM surrogate response plotted for comparison (third and final first-order dataset)

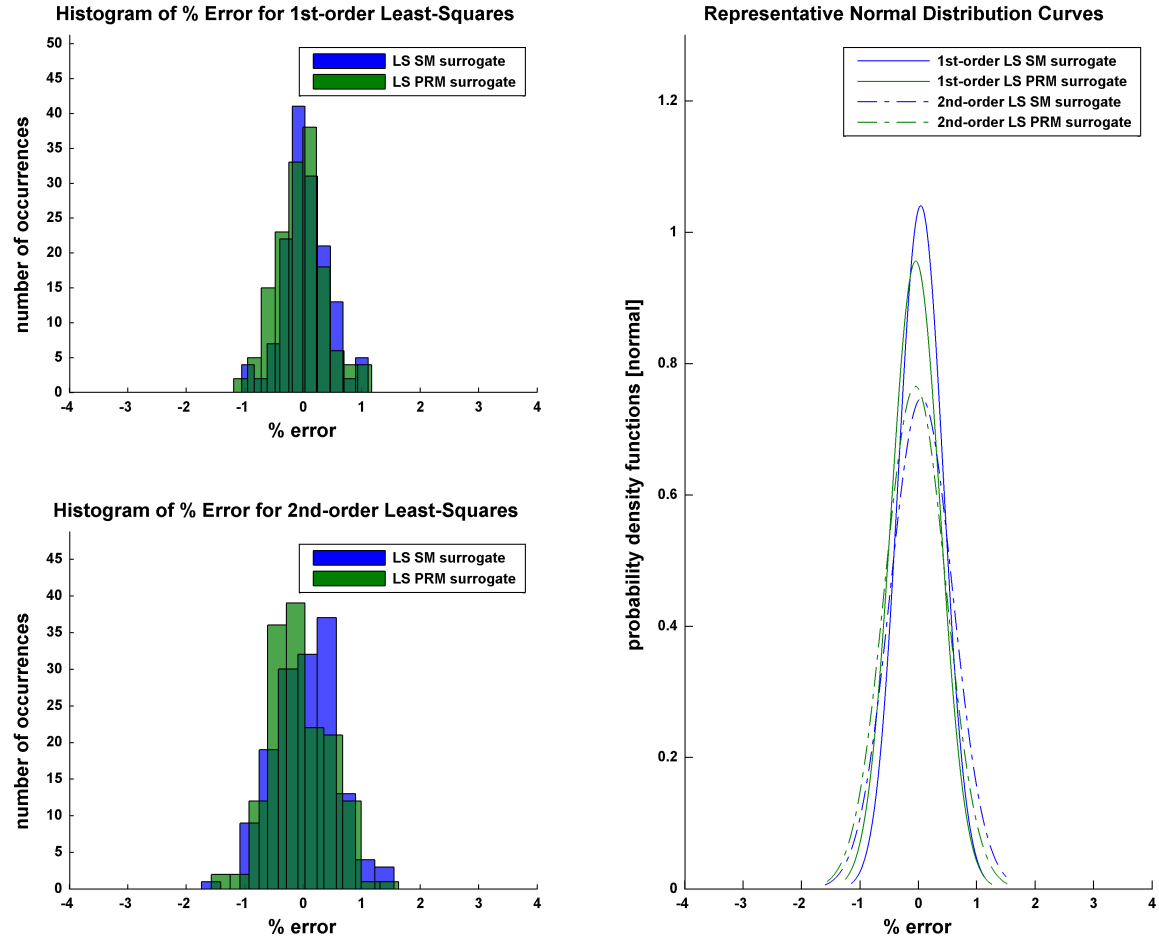


Figure 60. Histograms and representative normal distribution curves for the first and second-order least-squares surrogate models based on the second of three second-order datasets

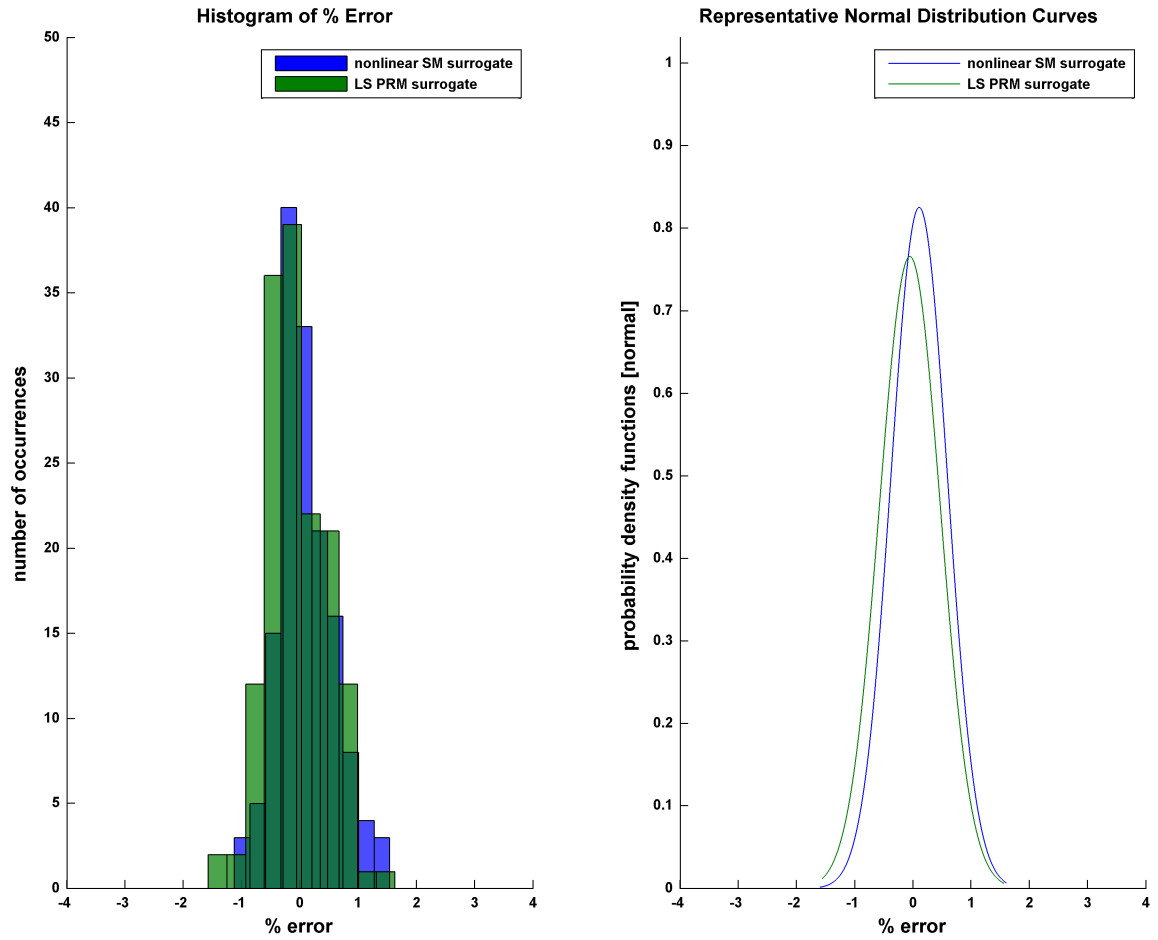


Figure 61. Histogram and representative normal distribution curve for the nonlinear surrogate model based on the second of three second-order datasets (second-order LS PRM surrogate data plotted for comparison)

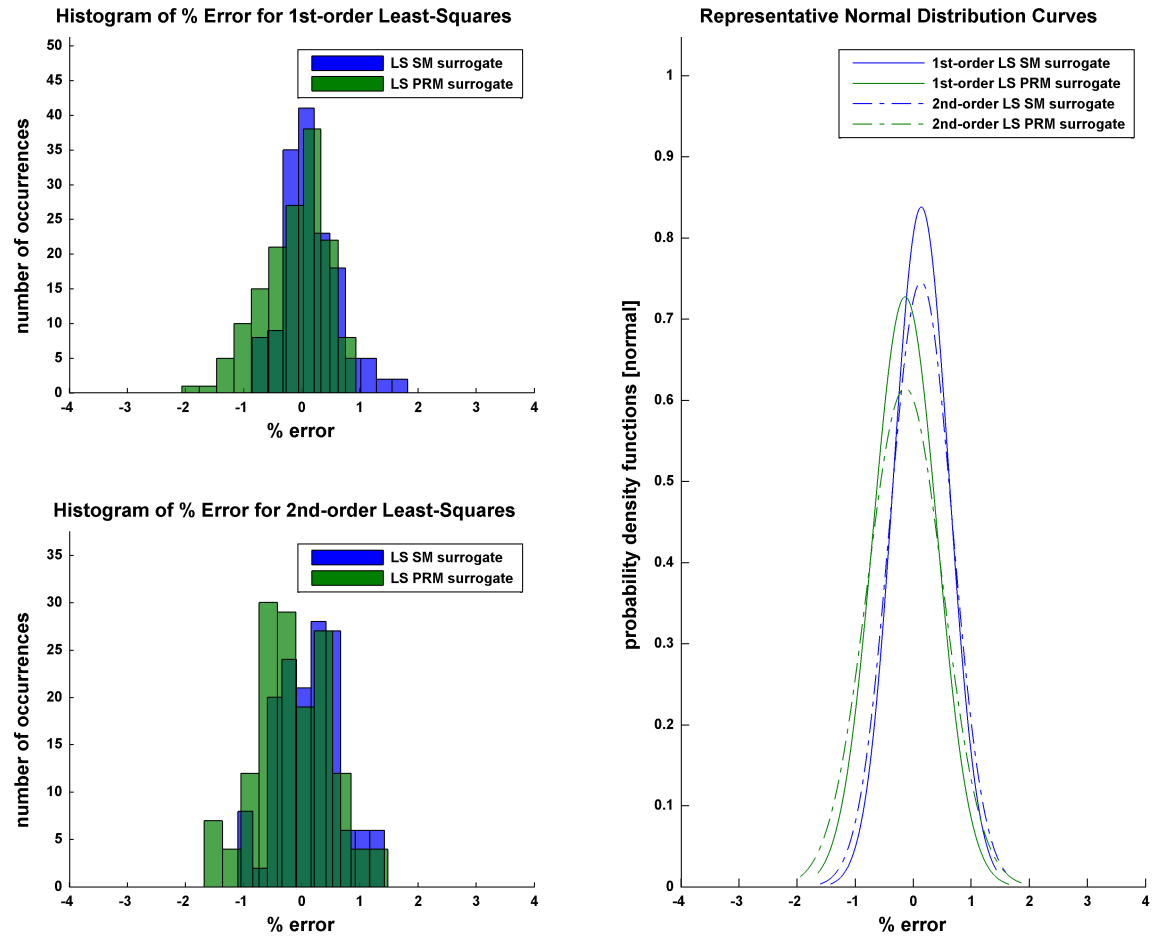


Figure 62. Histograms and representative normal distribution curves for the first and second-order least-squares surrogate models based on the third of three second-order datasets

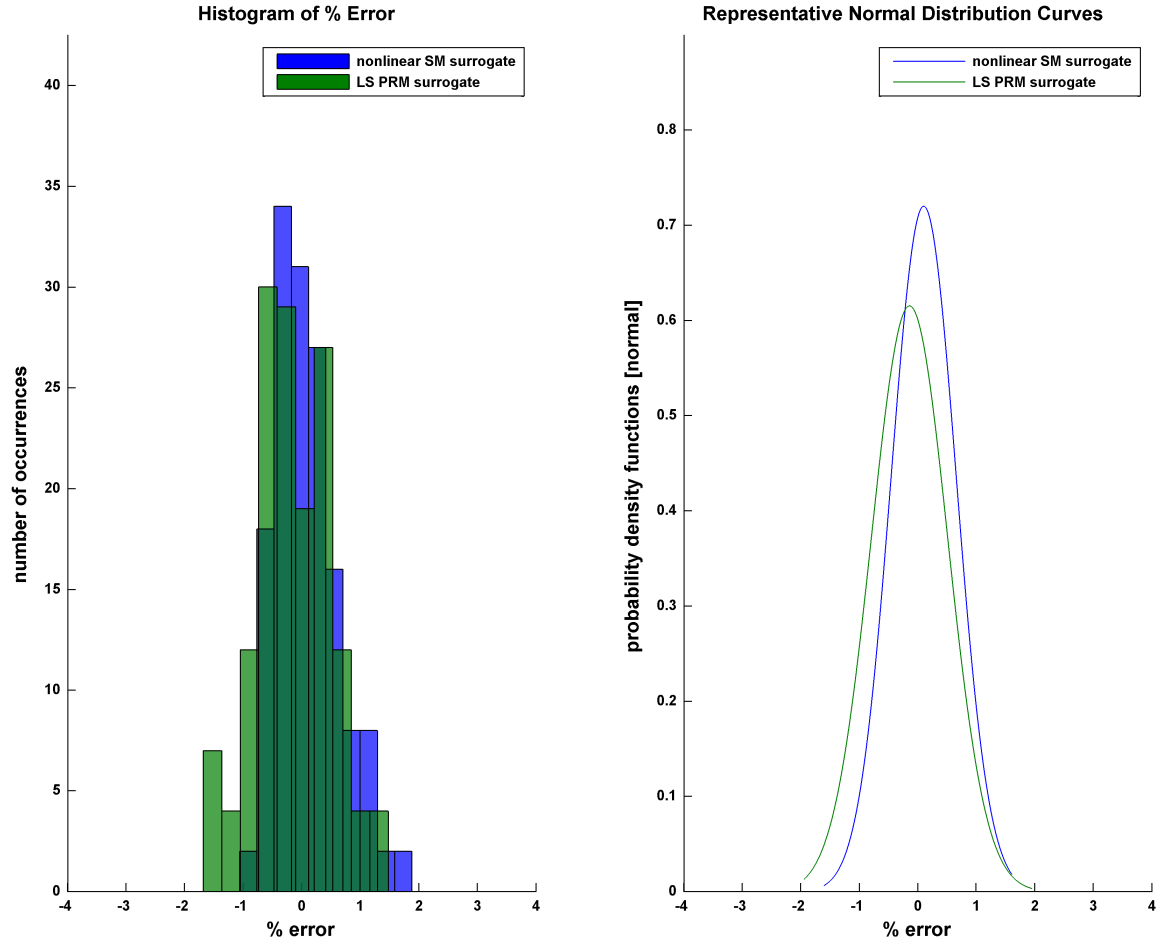


Figure 63. Histogram and representative normal distribution curve for the nonlinear surrogate model based on the third second-order dataset (second-order LS PRM surrogate data plotted for comparison)

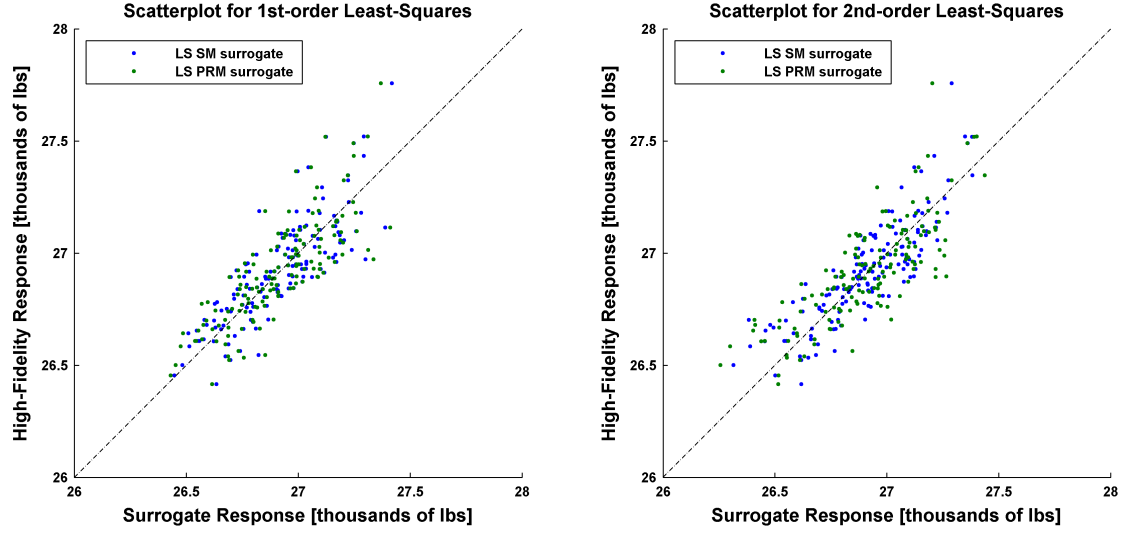


Figure 64. Scatterplot showing the surrogate model predictions against the true high-fidelity response for the first and second-order LS surrogate models for the first of three datasets

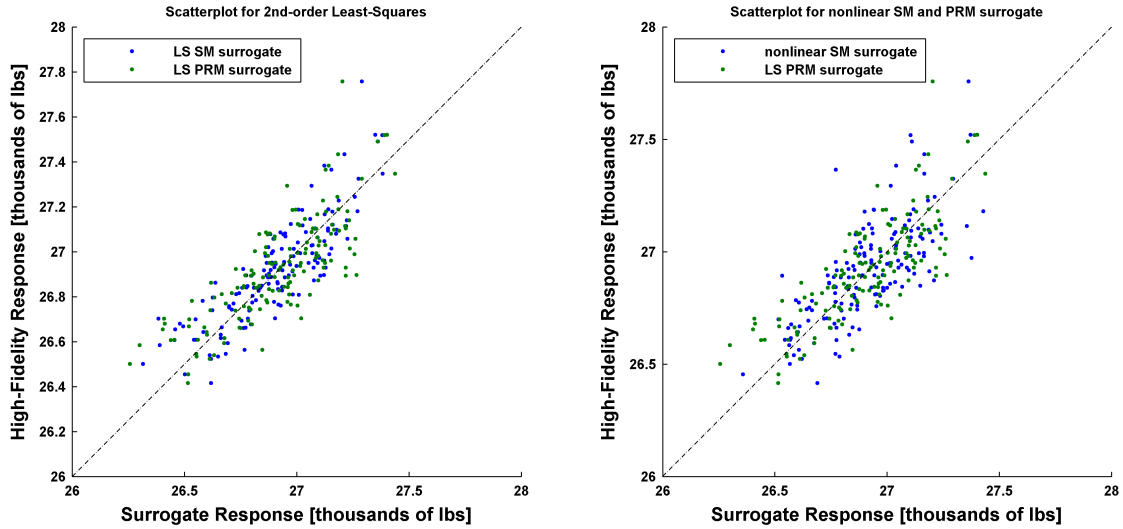


Figure 65. Scatterplot showing the surrogate model predictions against the true high-fidelity response for the second-order LS surrogate models as well as the non-linear polynomial-based SM surrogate for the first of three datasets

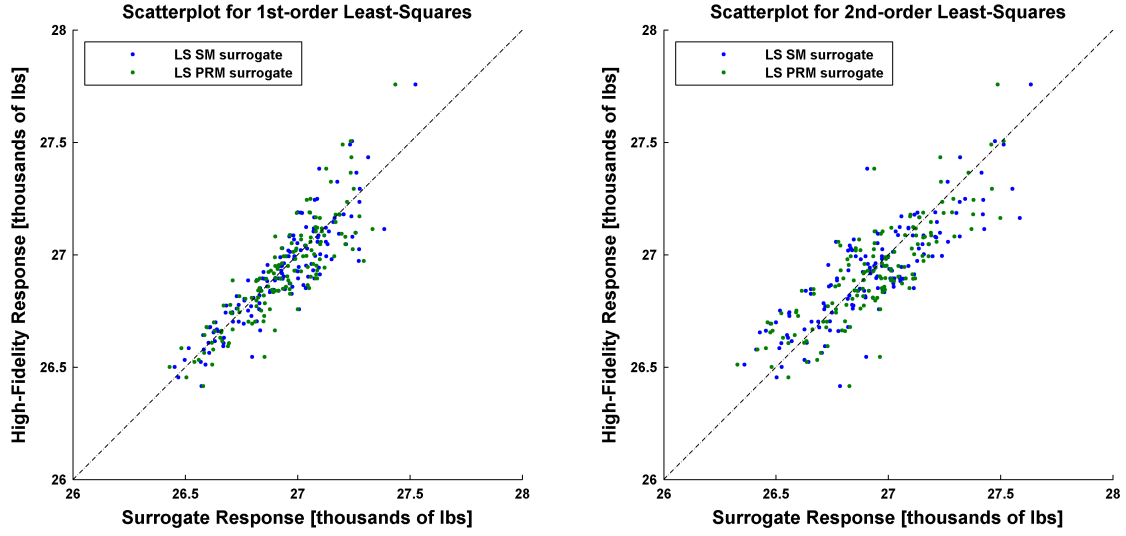


Figure 66. Scatterplot showing the surrogate model predictions against the true high-fidelity response for the first and second-order LS surrogate models for the second of three datasets

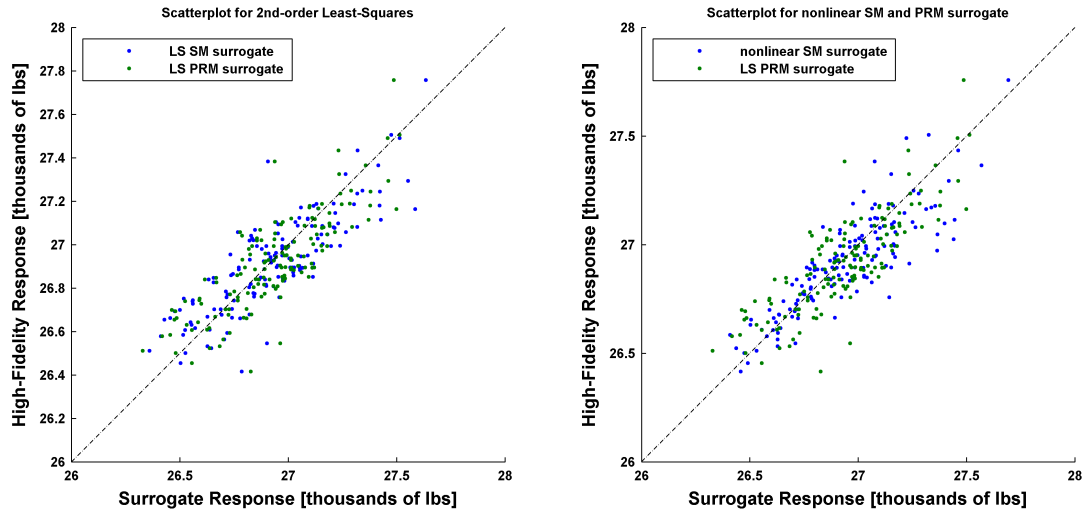


Figure 67. Scatterplot showing the surrogate model predictions against the true high-fidelity response for the second-order LS surrogate models as well as the non-linear polynomial-based SM surrogate for the second of three datasets

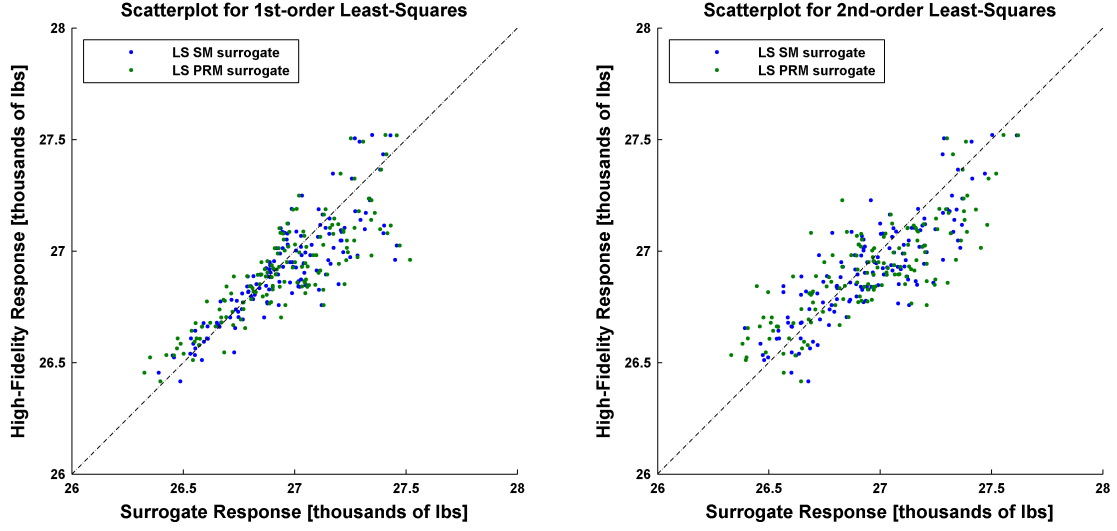


Figure 68. Scatterplot showing the surrogate model predictions against the true high-fidelity response for the first and second-order LS surrogate models for the third dataset

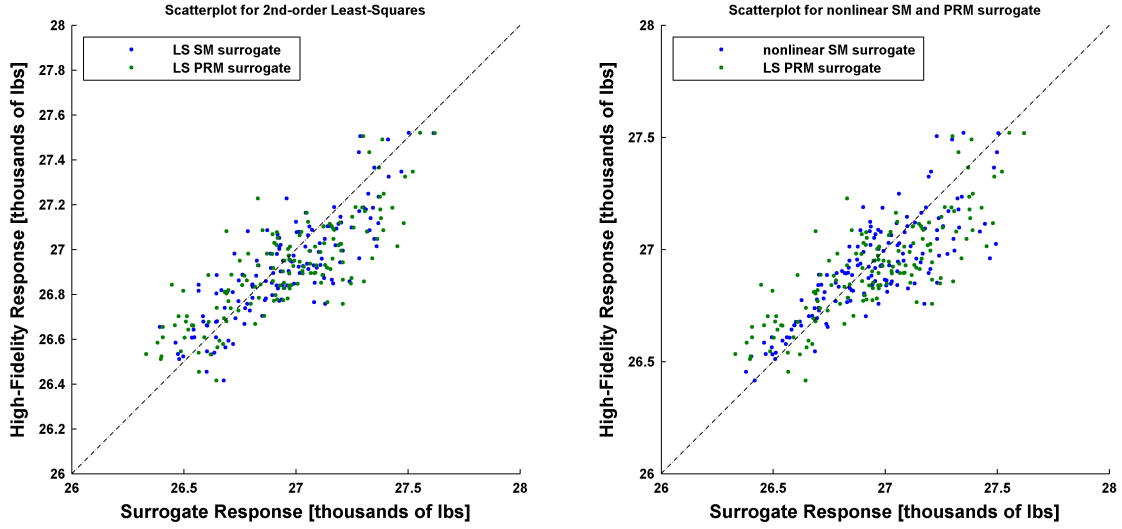


Figure 69. Scatterplot showing the surrogate model predictions against the true high-fidelity response for the second-order LS surrogate models as well as the non-linear polynomial-based SM surrogate for the third dataset

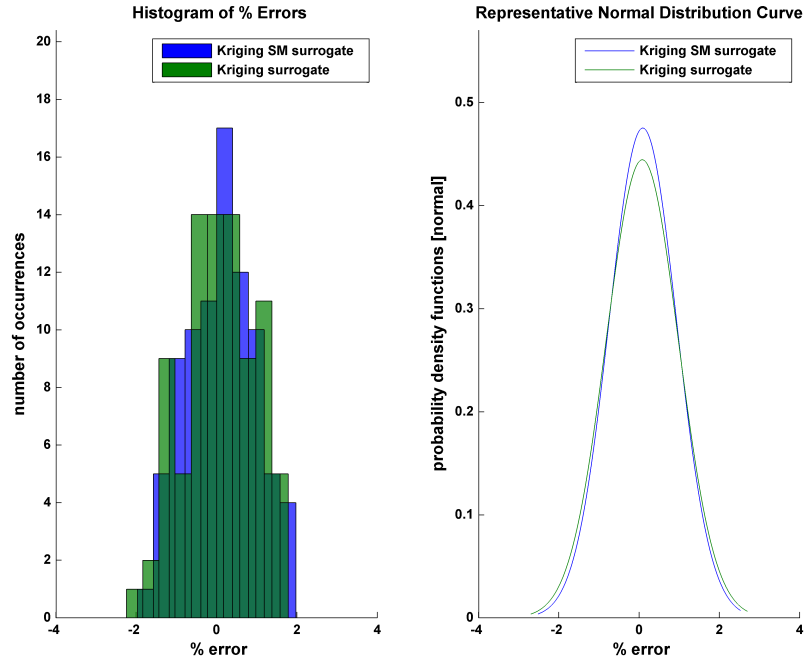


Figure 70. Scatterplot showing the surrogate model predictions against the true high-fidelity response for the kriging-based SM surrogate and the traditional kriging surrogate based on the same 58 sample data points from the first kriging dataset

References

1. Forrester, A., Sóbester, A., and Keane, A., *Engineering Design via Surrogate Modelling*, John Wiley and Sons, Inc., 2008.
2. Schaaf, W. L., *Carl Friedrich Gauss: Prince of Mathematicians*, The Moffa Press, Inc., 1964.
3. Strang, G., *Linear Algebra and Its Applications*, Brook/Cole Cengage Learning, 2006.
4. Myers, R. H. and Montgomery, D. C., *Response Surface Methodology: Process and Product Optimization Using Designed Experiments*, A Wiley-Interscience Publication, 2002.
5. Goel, T., Hafkta, R., and Shyy, W., “Comparing Error Estimation Measures for Polynomial and Kriging Approximation of Noise-free Functions,” *Structural and Multidisciplinary Optimization*, 2009.
6. Bohling, G., “Kriging,” *C&PE 940*, 2005.
7. Goovaerts, P., *Geostatistics for Natural Resources Evaluation*, Oxford University Press, 1997.
8. Lophaven, S., Nielsen, H., and Søndergaard, J., “DACE: A Matlab Kriging Toolbox,” *Informatics and Mathematical Modelling*, 2002.
9. Bandler and Cheng, “Space-Mapping: State of the Art,” *IEEE Trans. on Microwave Theory and Techniques*, Vol. 52:1, 2004.
10. Castro, Gray, and Guinta, “Sandia Report - Developing a Computationally Efficient Dynamic Multilevel Hybrid Optimization Scheme using Multifidelity Model Interactions,” Tech. Rep. SAND2005-7498, Sandia National Laboratories, 2005.
11. Arora, J. S., *Introduction to Optimum Design*, Elsevier, 2012.
12. Alyanak, E., Kolonay, R., Flick, P., Lindsley, N., and Burton, S., “Efficient Supersonic Air Vehicle Preliminary Conceptual Multi-Disciplinary Design Optimization Results,” *12th AIAA Aviation Technology, Integration, and Operations (ATIO) Conference*, 2012.
13. Alyanak, E. J. and Kolonay, R. M., “Efficient Supersonic Air Vehicle Structural Modeling for Conceptual Design,” *AIAA paper, public release 88ABW-2012-4638*, 2012.
14. Burton, S., Alyanak, E., and Kolonay, R., “Efficient Supersonic Air Vehicle Analysis and Optimization using SORCER,” *12th AIAA Aviation Technology, Integration, and Operations (ATIO) Conference*, 2012.
15. Niell, D. J., Herendeen, D. L., and Hoesly, R. L., “ASTROS Enhancements - Volume II - ASTROS Programmer’s Manual,” Tech. Rep. WL-TR-96-3005, Aerospace Vehicles Directorate (formerly the Flight Dynamics Directorate), 1995.
16. Raymer, D. P., *Aircraft Design: A Conceptual Approach (5th Ed.)*, American Institute of Aeronautics and Astronautics, Inc., 2012.

REPORT DOCUMENTATION PAGE			<i>Form Approved</i> <i>OMB No. 0704-0188</i>	
The public reporting burden for this collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden to Department of Defense, Washington Headquarters Services, Directorate for Information Operations and Reports (0704-0188), 1215 Jefferson Davis Highway, Suite 1204, Arlington, VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to any penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number. PLEASE DO NOT RETURN YOUR FORM TO THE ABOVE ADDRESS.				
1. REPORT DATE (DD-MM-YYYY) 27-03-14		2. REPORT TYPE Master's Thesis		3. DATES COVERED (From — To) 09-01-2012 – 27-03-2014
4. TITLE AND SUBTITLE A Method of Surrogate Model Construction which Leverages Lower-fidelity Information using Space Mapping Techniques			5a. CONTRACT NUMBER	
			5b. GRANT NUMBER	
			5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S) Thomas, Jason W., Captain, USAF			5d. PROJECT NUMBER	
			5e. TASK NUMBER	
			5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Air Force Institute of Technology Graduate School of Engineering and Management (AFIT/EN) 2950 Hobson Way WPAFB OH 45433-7765			8. PERFORMING ORGANIZATION REPORT NUMBER AFIT-ENY-14-M-46	
9. SPONSORING / MONITORING AGENCY NAME(S) AND ADDRESS(ES) AFRL/RQVC 2130 Eighth Street, Bldg 45, Rm 190 Wright-Patterson AFB OH 45433-7542 Alyanak, Edward J. Civ USAF AFMC AFRL/RQVC edward.alyanak.1@us.af.mil			10. SPONSOR/MONITOR'S ACRONYM(S) AFRL/RQ	
			11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION / AVAILABILITY STATEMENT Distribution Statement A. Approved for Public Release; Distribution Unlimited				
13. SUPPLEMENTARY NOTES This material is declared a work of the U.S. Government and is not subject to copyright protection in the United States				
14. ABSTRACT A new method of surrogate construction is developed and applied to a pair of computational tools used in the field of aircraft design. This new method involves the pairing of data sampled from the analytical model of interest with the execution of a similar analysis performed at a lower level of fidelity. This pairing is accomplished through the use of a space mapping technique, which is a process where the design space of a lower fidelity model is aligned a higher fidelity model. The intent of applying space mapping techniques to the field of surrogate construction is to leverage the information about a system's performance present at a lower fidelity level to bolster the predictive accuracy of a surrogate model based upon sampled data at a higher fidelity level. The results from the pairing of computational tools used in this research show modest gains in predictive accuracy for many of the cases investigated when compared to existing surrogate methodologies.				
15. SUBJECT TERMS multifidelity, surrogate, space mapping, surrogate method				
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT UU	18. NUMBER OF PAGES 111
a. REPORT U	b. ABSTRACT U	c. THIS PAGE U		
			19b. TELEPHONE NUMBER (Include Area Code) (937) (937) 255-3636 x4667 jeremy.agte@afit.edu	