

LOCAL CALIBRATION OF PAVEMENT-ME PERFORMANCE MODELS USING
MAXIMUM LIKELIHOOD ESTIMATION

By

Rahul Raj Singh

A DISSERTATION

Submitted to
Michigan State University
in partial fulfillment of the requirements
for the degree of

Civil Engineering – Doctor of Philosophy

2024

ABSTRACT

The mechanistic empirical pavement design guide (MEPDG) is a state-of-the-art design approach that incorporates material properties, traffic, and climate to estimate the incremental damage using mechanical responses of the pavement. The cumulative damage is used to predict the field distress using empirical transfer functions. The Pavement-ME transfer functions have been nationally calibrated using long-term pavement performance (LTPP) pavement sections and other experimental test section data such as MnRoad. These nationally calibrated models may not represent the construction practices, materials, and climatic conditions of a particular state/region. Studies have calibrated the Pavement-ME transfer functions using the least squares method. Least squares is a widely used simplistic method based on the normal independent and identically distributed (NIID) assumption. Literature shows that these assumptions may not apply to non-normal distributions. This study introduces a new methodology for calibrating the bottom-up cracking, total rutting, and international roughness index (IRI) models in new flexible pavements and the transverse cracking and IRI models in new rigid pavements using Maximum Likelihood Estimation (MLE). The approach in this study includes MLE using synthetic and observed data, and the results are compared with those of the least squares approach. The MLE and least squares methods were also combined with resampling techniques to improve the robustness of calibration coefficients. The data are analyzed from the Michigan Department of Transportation's (MDOT) Pavement Management System (PMS) database to obtain the pavement sections and observed performance data for calibration.

Despite several calibration efforts, limited research is available on the impact of calibration on pavement design. The calibrated models using the least squares method were then used for pavement design to estimate the calibration effects and compare them with AASHTO93 designs. Based on the newly calibrated coefficients, 44 new flexible and 44 rigid sections were designed. This study also identifies the controlling distresses for pavement design.

It is often not viable to calibrate all coefficients at the same time. Therefore, it is crucial to identify the most sensitive transfer function coefficients. Moreover, the sensitivity also indicates the impact of each coefficient on the performance prediction. Typically, the sensitivity is obtained using a normalized sensitivity index (NSI). This study estimated the sensitivity of the Pavement-ME transfer function coefficients using scaled sensitivity coefficients (SSCs).

The results show that MLE outperformed the least squares method for non-normally distributed data, such as transverse cracking and bottom-up cracking models for synthetic and observed data. Using the calibrated models for pavement design showed that, on average, the surface thicknesses using locally calibrated coefficients were thinner by 0.22 and 0.44 inches for flexible and rigid pavements, respectively. Critical design distresses for flexible pavements include bottom-up and thermal cracking. On the other hand, transverse cracking and IRI control the designs for rigid sections. The sensitivity of Pavement-ME model coefficients showed that SSCs provide a more reliable sensitivity on a range of independent variables rather than a point estimate, unlike NSI. Overall, this study helps improve the calibration process for local conditions.

This dissertation is dedicated to my parents,
Mr. Dinesh Singh and Mrs. Ranjana Singh,
and my sister, Shweta Rani.

ACKNOWLEDGMENTS

First, I would like to extend my deepest gratitude to my advisor, Dr. Syed Waqar Haider, for his continuous support and guidance. His experience, thorough knowledge, patience, and involvement in every step of this project helped me complete this thesis successfully and on time.

I want to express my gratitude to Dr. M. Emin Kutay for his guidance throughout my Ph.D. journey. I was fortunate to have taken a course with Dr. Kirk D. Dolan, for which I am thankful to him. It helped me build a foundation in parameter estimation, publish papers, and implement it into my thesis. I would also like to thank Dr. Karim Chatti for his continuous feedback and for serving on my Ph.D. committee.

I want to thank my parents, Mr. Dinesh Singh and Mrs. Ranjana Singh, and my sister, Shweta Rani, for their immense love and support. I would also like to thank my friends Komal Narendrakumar Thakkar and Divyang Narendrakumar Thakkar for making my journey to the United States easier.

I have been blessed to have worked with some brilliant minds in the Department of Civil and Environmental Engineering at Michigan State University. I want to thank Hamad Muslim, Mumtahir Hasnat, Faizan Lali, Mahdi Ghazavi, Poornachandra Vaddy, Farhad Abdollahi, and many others for being part of my academic journey.

Finally, I would like to thank God for giving me the strength and this opportunity to be a part of such a prestigious institute and work with such kind and generous people.

TABLE OF CONTENTS

CHAPTER 1 - INTRODUCTION.....	1
CHAPTER 2 - LITERATURE REVIEW.....	6
CHAPTER 3 - DATA FOR CALIBRATION	36
CHAPTER 4 - METHODOLOGY	78
CHAPTER 5 - RESULTS AND DISCUSSION	107
CHAPTER 6 - CONCLUSIONS, RECOMMENDATIONS AND FUTURE SCOPE.....	137
REFERENCES	146
APPENDIX.....	151

CHAPTER 1 - INTRODUCTION

1.1 BACKGROUND

The AASHTOWare Pavement Mechanistic-Empirical Design (Pavement-ME) is the latest American Association of State Highway and Transportation Officials (AASHTO) pavement design software edition. It is based on the AASHTO's Mechanistic-Empirical Pavement Design Guide (MEPDG). Pavement-ME is a significant shift from the empirical design process developed and supported by the AASHTO Interim Guide for Design of Pavement Structures (AASHTO 1972) through the AASHTO Guide for Design of Pavement Structures and its 1998 Supplement (AASHTO 1998) (1). While these earlier AASHTO design guides are based on empirically derived performance equations developed using data from the AASHO road test conducted in the 1950s, these have been widely popular for pavement design. About 48 agencies reported using the AASHTO empirical design guides after their refinements provided by AASHTO in 1986 and 1993 (1). Despite the refinements in the material input parameters and the design reliability, the previous design guides' empirical nature limits their performance for the following reasons (2).

- The application of the AASHO Road test is limited by its specific geographic location, which does not account for the climatic effects of a different location on pavement performance.
- Truck traffic volume has increased significantly since the 1960s, and truck configurations have also changed.
- All test sections were built using a single hot mix asphalt (HMA) mixture for flexible pavements and one Portland cement concrete (PCC) mixture for rigid pavements over one subgrade soil type.

Recognizing these limitations, the Joint Task Force on Pavements (JTFP) initiated an effort in 1996 to develop the MEPDG using mechanistic pavement design principles. The new mechanistic-empirical (M-E) design procedure offers multiple benefits, taking advantage of improvements in material characterization, axle load spectra, and climate models to predict the pavement's performance.

Version 0.7, the research version of the MEPDG software, was first released in July 2004. It was revised several times under different projects funded by the National Cooperative Highway Research Program (NCHRP). The software's revisions included the release of version 0.8 in November 2005, version 0.9 in July 2006, version 1.0 in April 2007, and version 1.1 in September 2009. An MEPDG Manual of Practice was published in 2008, aiming to assist highway agencies in implementing the M-E design method with version 1.0. It was adopted as an interim AASHTO pavement design procedure a year earlier in 2007 (3). Another version of the M-E design was released in April 2011 called Design, Analysis, and Rehabilitation for Windows (DARWin). DARWin-ME software was later named AASHTOWare Pavement METM once AASHTO underwent rebranding in 2013. Currently, the latest version of Pavement-ME software is version 2.6.2, and the online version is version 3.0. In addition, a Backcalculation Tool (BcT) and Calibration Assistance Tool (CAT) have been developed for use with the Pavement-ME software.

1.2 PROBLEM STATEMENT

The MEPDG was developed under the NCHRP project 1-37A (4) to overcome the limitations of the AASHTO 1993 method (5). It is an advanced pavement design tool for new and rehabilitated pavements. MEPDG incorporates material properties, traffic, and climate to estimate the incremental damage using mechanical responses of the pavement. The cumulative damage is empirically used to predict the field distress using transfer functions. The transfer functions used in the Pavement-ME have been globally calibrated using the Long-term Pavement Performance (LTPP) pavement sections (6). Although the globally calibrated models provide fair performance predictions for the entire US road network, these may not represent the construction practices, materials, and climatic conditions of a particular state/region. Therefore, nationally calibrated models may underpredict or overpredict the pavement performance in specific states or regions. Recalibration of these models has been recommended for local conditions in the local calibration guide (7). The design distresses in the Pavement-ME include transverse cracking (percentage of slabs cracked), transverse joint faulting (inches), and international roughness index (IRI in inches/mile) for rigid pavements. For flexible pavements, the design distress includes bottom-up cracking (percentage), top-down cracking (percentage), rutting (inches), thermal (transverse) cracking (feet/mile), reflective cracking (feet/mile), and IRI (inches/mile).

Several studies have been performed in Michigan in the recent past to characterize climate, traffic, and material properties, as well as to calibrate the performance models to address the local conditions, materials, and construction practices in the Pavement-ME procedure (8-10). While all the material properties and calibration of performance models were addressed to improve the Pavement-ME local applicability and accuracy, there were still some data gaps, specifically for material characterization and pavement construction. Examples of past data gaps include clustered traffic data, HMA mix, and binder properties. Gaps in data need to be estimated (corresponds to Level 3 for Pavement-ME input levels), which may not be accurate for the location; therefore, having the actual values for new projects will likely improve Pavement-ME calibration accuracy. Also, a limited number of rigid pavement sections were available for previous Michigan calibration efforts; therefore, adding more data from new sections would improve the performance model prediction.

Most calibration studies have used the least squares approach to calibrate the Pavement-ME transfer functions. Least squares is a widely used simplistic method based on the normal independent and identically distributed (NIID) assumption. The NIID assumption states that observations in a sample are independent, i.e., the occurrence of one does not influence another. Additionally, these observations should have identical probability distributions, i.e., drawn from the same underlying population distribution. Furthermore, the assumption implies that the observed data and error term follow a normal distribution. Literature shows that the least squares method assumptions may not apply to the non-normal distributions. This limits the robustness of the least squares method for transverse cracking in rigid pavements and bottom-up cracking in flexible pavements, which are usually non-normally distributed.

The ultimate goal of Pavement-ME calibration is improving pavement designs for local conditions. Despite several calibration efforts, limited research is available on the effect of calibration on pavement design. Estimating the change in design thicknesses and identifying critical distresses using the calibrated models is vital. By understanding which distress types are most relevant to a region, agencies can develop mitigation and maintenance strategies leading to longer pavement service lives.

State Highway Agencies (SHAs) often struggle to identify the most critical data collection needs since the Pavement-ME requires several design inputs. Several studies have conducted sensitivity analyses to determine the most sensitive inputs to the distress prediction

models for new and rehabilitated pavements to address this issue. However, limited research is available to assess the impact of each calibration coefficient on the predicted pavement distress and performance. These studies quantified the sensitivity of coefficients using a sensitivity index and a typical range of design inputs. The sensitivity metric adopted to accomplish the sensitivity analyses is called the normalized sensitivity index (NSI), defined as the percentage change of predicted distress relative to its global prediction caused by a given percentage change in the coefficient. While NSI can rank the coefficients based on their level of sensitivity, it does not provide information about any potential correlation between them or how accurately these can be estimated. Moreover, since the calculation of NSI requires distress data, its magnitude can change if the data source is changed; hence, the sensitivity ranking of the coefficients may vary, as reported by Dong et al. (11).

1.3 RESEARCH OBJECTIVES

The recalibration of the Pavement-ME models is crucial for any SHA implementing M-E design. This includes identifying the suitable Pavement-ME inputs, potential projects, and performance data. It is also essential to verify the feasibility of the calibrated models for pavement design. The main objectives of this study are to (a) calibrate the Pavement-ME models using improved inputs (traffic, HMA and climate) and additional data (potential projects and performance data) for new flexible and rigid pavements, (b) assess the impact of calibrated models on design thicknesses and to identify critical design distresses, (c) apply maximum likelihood estimation (MLE) to calibrate and validate the Pavement-ME models and compare the results with the least squares method, (d) determine the sensitivity of Pavement-ME calibration coefficients over a continuous scale of independent variables using scaled sensitivity coefficients (SSCs) and compare it with the traditional NSI approach.

These objectives were accomplished using the pavements and the corresponding performance data from the MDOT Pavement Management System (PMS) database.

1.4 DISSERTATION OUTLINE

This dissertation contains six chapters. Chapter 1 outlines the background of the Pavement-ME, the problem statement, and the research objectives. Chapter 2 documents the literature review from previous calibration studies, Pavement-ME transfer functions, and calibration approaches.

Chapter 3 discusses the input and performance data used for calibration efforts. This includes data collection efforts, a summary of the performance, and input data for the selected pavement sections for model calibrations. Chapter 4 details the local calibration methods and procedures used in this study. This chapter also includes the methodology used for calculating the SSCs. Chapter 5 presents the local calibration results for the various performance prediction models, including calibration results from the least squares and MLE methods. This chapter also consists of the results from assessing the impact of calibration on pavement design and the SSC plots. Chapter 6 summarizes this study's conclusions, recommendations, and future scope. Each chapter has a summary at the end, which outlines the overall highlights of the chapter.

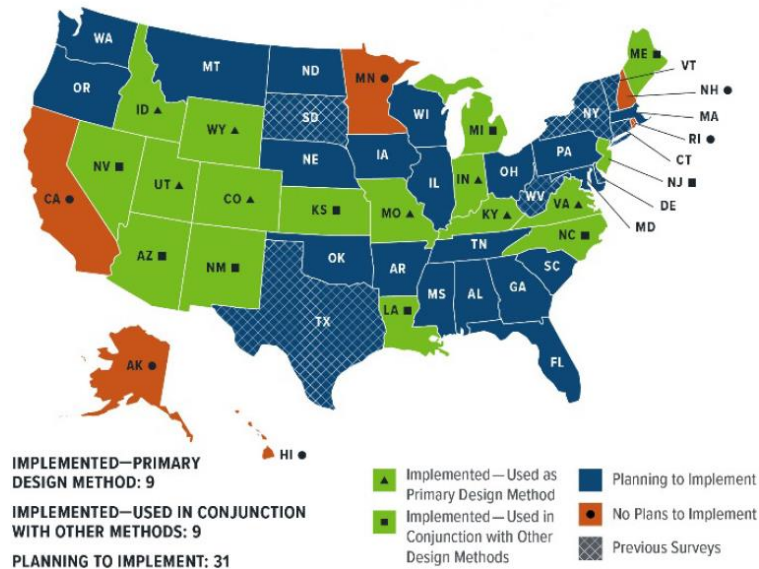
CHAPTER 2 - LITERATURE REVIEW

The Pavement-ME provides highway agencies with a practical tool for designing new and rehabilitated pavements. The analyses in M-E principles use primary pavement responses (stresses, strains, and deflections) and incremental damage over time to predict surface distress through transfer functions. The reliability of performance prediction models depends on the accuracy of the transfer functions, which is achieved through calibration and subsequent validation with observed pavement condition data. A satisfactory correlation between measured and predicted performance indicators increases the viability, acceptance, and usage of the MEPDG procedures for pavement analysis and design procedures. Calibration is a mathematical procedure to reduce the difference between predicted and measured distress values. Validation refers to a process that evaluates the performance of mathematical models on an independent dataset (i.e., data not used for model development). This chapter outlines the literature review of calibration approaches, the methodology used in different studies, and the concept of reliability for Pavement-ME predictions.

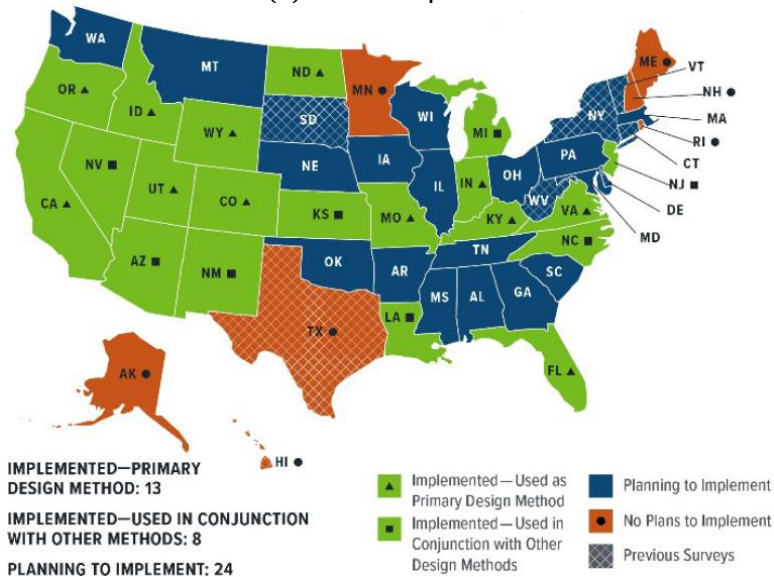
2.1 IMPLEMENTATION OF PAVEMENT-ME

The AASHTO93 empirical pavement design method has been popular and used by highway agencies for several decades (5). Highway agencies are still using it as their current pavement design procedure. The shift from an empirical to a more M-E design method occurred in 2008 after the publication of the MEPDG practice manual and the release of Pavement-ME software (3). The adoption of the Pavement-ME design was further enhanced by publishing the local calibration guide to implement nationally calibrated models for local conditions (7). In recent years, other supplemental tools like the Calibration Assistance Tool (CAT) and Backcalculation Tool (BcT) have helped agencies implement the Pavement-ME design. The adoption of Pavement-ME design started soon after its release, with fifteen state highway agencies (SHA) implementing it within the first few years (1). The implementation significantly increased between 2010 and 2020 and became stagnant due to several challenges. These challenges include the unavailability of input data, pavement sections for calibration, and sufficient good-quality performance data. Some agencies have returned to using their original design practice (usually AASHTO-93) or M-E design in parallel with their original method. As of 2021, nine state

agencies are using Pavement-ME as their primary design method for flexible pavements, and thirteen are using it for rigid pavements. Further, nine state agencies use Pavement-ME with other design methods for flexible pavements, whereas eight use it for rigid pavements (12). Figure 2-1 shows the implementation status of the Pavement-ME design for flexible and rigid pavements.



(a) Flexible pavements



(b) Rigid pavements

Figure 2-1 Pavement-ME implementation status (12)

2.2 LOCAL CALIBRATION EFFORTS

Calibration of Pavement-ME models is an optimization problem. Several researchers have calibrated these models using different optimization methods. This section summarizes the calibration methods and efforts in different states.

2.2.1 Least Squares Method

The least squares method is a mathematical technique used to minimize the sum of squared differences between observed and predicted values. Calibration of the Pavement-ME transfer functions is established by minimizing the bias and standard error between the measured and predicted distress. Researchers have used several simplistic and robust approaches leveraging the least squares method for calibration. Hall et al. (2011) used the Microsoft Excel solver function to calibrate the alligator cracking model for flexible pavements in Arkansas (13). Tarefder and Rodriguez-Ruiz (2013) calibrated the rutting, alligator cracking, and longitudinal cracking models for flexible pavements in New Mexico. The process involved changing the calibration coefficients and rerunning Pavement-ME in an iterative process to obtain minimum mean residual error (MRE) and the sum of squared errors (SSE) (14). These calibration efforts have become more robust with the development of computational and statistical techniques. Dong et al. (2020) calibrated the joint faulting model for rigid pavements in Ontario. This study used three different optimization techniques: (1) one at a time using trial and error; (2) Microsoft Excel solver function; (3) Levenberg-Marquardt Algorithm (LMA). Results showed that calibration using approaches (2) and (3) significantly improved the bias and standard error of estimate (SEE) (11). Haider et al. (2020) calibrated the transverse cracking and IRI models for rigid pavements in Michigan. This study used resampling methods like bootstrapping and repeated split sampling for calibration and validation. The results showed that resampling methods provide a more robust calibration than traditional methods, along with confidence intervals of the SEE, bias, and transfer function coefficients (9). Tabesh and Sakhaeifar (2021) calibrated rutting, IRI, top-down, and bottom-up cracking models in Oklahoma using a narrow-down iterative approach in Microsoft Excel solver (15). This study showed significant improvement in the Pavement-ME predictions and flexible pavement designs. All these studies have used the least squares to calibrate these transfer functions using the NIID assumption. Although least squares is a popular and simplistic approach, the assumptions may not be valid,

especially for non-normally distributed data. Tables 2-1 and 2-2 summarize the calibration efforts from different states.

Table 2-1 Summary of calibration efforts for flexible pavements

States	Number of sections		Pavement-ME models	Version	Year
	New	Rehabilitation			
Arkansas (16)	38	-	BU, TD, RUT, TC,	V1.1	2014
Colorado (17)	46	49	BU, RUT, TC, IRI	V1.0	2013
Minnesota (18)	39		BU, RUT, TC	V1.0	2009
Montana (19)	102		RUT, TC	V0.9	2007
New Mexico (14)	19	5	BU, TD, RUT, IRI	V1.0	2013
Ohio (20)	13	-	RUT, IRI	V1.0	2009
Oregon (21)	-	38	RUT, BU, TD, TC	V2.0.19	2019
South Carolina (22)	14	-	RUT, BU, TD	V 2.2	2016
Utah (23)	21	9	RUT	V 1.0	2009
Washington (24)	8	-	BU, TD, RUT	V 1.0	2009
Arizona (25)	58	42	BU, TD, RUT, IRI, REF	DARWin- ME	2014
Iowa (26)	35	-	BU, TD, RUT, IRI	V1.1	2014
Kansas (27)	28	-	TD, RUT, IRI	-	2015
Michigan (28)	163	121	BU, TD, RUT, TC, IRI, REF	V2.6	2023
North Carolina (29)	46	-	BU, RUT	DARWin- ME	2011
Texas (30)	18	-	RUT	-	2009
Wyoming (31)	86	-	BU, RUT	V2.2	2015
Missouri (32)	6	11	BU, TD, RUT, TC, IRI, REF	V2.5.5	2020
Georgia (33)	27	20	BU, RUT, TC	-	2014
Louisiana (34)	71	33	BU, RUT, REF	V2.0	2016
Virginia (35)	53	59	BU, RUT, IRI	V2.2.6	2022
Tennessee (36)	-	76	BU, TD, RUT, IRI	V2.1	2016
Oklahoma (15)	65	-	BU, TD, RUT, TC, IRI	V2.3	2021

Note: BU = Bottom-up cracking; TD = Top-down cracking; TC = Thermal cracking; RUT = Total rutting; REF = Reflective cracking; IRI = International roughness index

Table 2-2 Summary of calibration efforts for rigid pavements

States	Number of sections		Pavement-ME models	Version	Year
	New	Rehabilitation			
Colorado (17)	25	7	TC, JF, IRI	V1.0	2013
Minnesota (18)	65		TC	V1.0	2009
Ohio (20)	14	-	IRI	V1.0	2009
South Carolina (22)	6	-	TC	V 2.2	2016
Arizona (25)	48	-	TC, JF, IRI	-	2014
Kansas (27)	32	-	JF, IRI	V1.3	2015
Michigan (28)	46	11	TC, JF, IRI	V2.6	2023
Wyoming (31)	26	-	JF	V2.2	2015
Missouri (32)	33	9	TC, JF, IRI	V2.5.5	2020
Georgia (33)	9	2	TC, JF	-	2014
Louisiana (34)	43	-	TC, JF	V2.0	2016
Idaho (37)	40	-	TC, JF, IRI	V2.5.3	2019
Virginia (35)	17	-	JF, IRI	V1.3	2022

Note: TC = Transverse cracking; JF = Joint faulting; IRI = International roughness index

2.2.2 Maximum Likelihood Estimation (MLE) Method

MLE has been used by several researchers in different fields; limited research is available on the use of MLE to calibrate the Pavement-ME transfer functions. Chen et al. (2021) presented a local calibration model for predicting punchout distress in continuously reinforced concrete pavement (CRCP). This study utilized a Weibull distribution to estimate the number of equivalent single axle loads (ESALs) leading to punchout, employing MLE and a Newton method. The model was validated using data from the LTPP database, demonstrating its efficacy in describing punchout behavior and facilitating predictions for CRCP reliability and rehabilitation planning (38). Haider et al. (2023) showed the robustness of MLE for non-normally distributed data using the MDOT PMS database. The bias for the transverse cracking model in rigid pavements and the bottom-up cracking model in flexible pavements was significantly improved (28).

MLE stands out as an advantageous and robust method for parameter estimation as it is based on a well-defined likelihood function rooted in the underlying probability distribution of the data. MLE is computationally efficient, leveraging standard probability distributions, making it usable for multi-dimensional and complex models. MLE excels in estimating parameters for probabilistic models, and it is especially useful in machine learning (39). Unlike the least squares method, MLE shows resilience to outliers as the probability of outliers is very low and offers a potential advantage in the bias-variance tradeoff. The bias-variance tradeoff is used in statistical modeling and machine learning to balance between capturing the underlying pattern in the data (low bias) and resisting sensitivity to fluctuations and noise (high variance). Models with high bias oversimplify data, leading to underfitting, while those with high variance overfit and fail to generalize the model for new data (40). The bias-variance tradeoff highlights the importance of finding the optimal model complexity and employing regularization or ensemble methods to strike the right balance. Understanding this tradeoff is crucial for effective model selection and evaluation, emphasizing the need for ample high-quality training data to minimize bias and variance in overall error.

Jose (2023) showed the application of MLE in modeling commodity prices and pricing financial derivatives. This study highlighted estimating model parameters using various methods, with a preference for maximum likelihood when the parametric specification is highly trusted. The comparison in the study evaluates different techniques for obtaining maximum likelihood estimates in the context of Ornstein-Uhlenbeck mean-reverting models based on observations

collected at arbitrary points in time (41). Pan and Fang (2002) discussed MLEs for parameters in growth curve models, emphasizing their differences from generalized least squares estimates (GLSE). The special case of Rao's simple covariance structure (SCS), where MLEs coincide with GLSEs, facilitating analytical and tractable statistical inferences in growth curve models, is explored. It also delves into the restricted maximum likelihood (REML) estimate under the assumption of the SCS, offering insights into statistical techniques for analyzing growth curve models (42). Myung (2003) illustrated using MLE, stressing its fundamental role in statistical inference. Moreover, this study emphasized using MLE and its superiority in nonlinear modeling with non-normal data (39). Bauke (2007) showed the limitations of using the least squares method for estimating power-law distribution exponents due to incompatible assumptions with empirical data. It shows the advantages of maximum likelihood estimators, deemed reliable for power-law distributions, with asymptotic efficiency (43). Zhang and Callan (2001) addressed the information filtering systems based on statistical retrieval models, focusing on optimizing dissemination thresholds for document delivery. This study introduced a novel algorithm grounded in the maximum likelihood principle to adjust thresholds by explicitly compensating for bias in relevant information obtained during filtering. Experiments using Text Retrieval Conference (TREC)-8 and TREC-9 filtering track data illustrate the algorithm's effectiveness in jointly estimating parameters and improving system performance. The TREC is an annual series of workshops evaluating information retrieval systems (44). Rayner and MacGillivray (2002) showed the use of numerical maximum likelihood estimation for distributions defined only by quantile functions, focusing on the g-and-k and generalized g-and-h distributions. Despite increased computing power, this aspect of MLE has received limited attention. This study presents and investigates numerical MLE procedures, conducts simulation studies, and emphasizes the need for resampling to obtain reliable estimates for quantile-defined distributions through maximum likelihood (45). Lio and Liu (2020) performed the regression analysis by defining a likelihood function using uncertain measures to represent parameter likelihoods. This study employs MLE for uncertain regression models, simultaneously calculating the uncertainty distribution of the disturbance term. Numerical examples demonstrate the proposed method, emphasizing its applicability to cases with imprecise observations. Future research directions include applying uncertain maximum likelihood to parameter estimation in uncertain differential equations, time series analysis, and hypothesis testing (46).

2.3 PAVEMENT-ME PERFORMANCE MODELS

The following section presents the formulation of transfer functions for flexible pavement models and the local calibration coefficients for different states.

2.3.1 Performance Models for Flexible Pavements

2.3.1.1. Fatigue cracking (bottom-up)

Bottom-up cracking is a load-related distress caused by the repeated axle load. These cracks initiate at the bottom of the asphalt concrete (AC) layer and propagate to the surface. The total cumulative damage DI can be estimated by summing the cumulative damage that is computed using Miner's law (47), as shown in Equation (2-1).

$$DI = \sum (\Delta DI)_{j,m,l,p,T} = \sum \left(\frac{n}{N_{f-HMA}} \right)_{j,m,l,p,T} \quad (2-1)$$

where,

n = Number of actual axle load applications within a specific time period

j = Axle load-interval

m = Axle type (single, tandem, tridem, quad)

l = Truck type classified in the MEPDG

p = Month

T = Median temperature for five temperature quintiles used in MEPDG

N_{f-HMA} = Allowable number of axle load applications, which can be computed using Equation (2-2).

$$N_{f-HMA} = C \times k_1 \times C_H \times \beta_{f1}(\varepsilon_t)^{-k_2\beta_{f2}}(E_{HMA})^{-k_3\beta_{f3}} \quad (2-2)$$

where,

ε_t = Tensile strain at critical AC locations

E_{HMA} = Dynamic modulus (E^*) of the Hot mix asphalt (HMA), psi

k_1, k_2, k_3 = Laboratory regression coefficients, and $\beta_{f1}, \beta_{f2}, \beta_{f3}$ = local or field calibration constants

C = Adjustment factor (laboratory to the field) as shown in Equation (2-3) and Equation (2-4).

$$C = 10^M \quad (2-3)$$

$$M = 4.84 \left(\frac{V_{be}}{V_a + V_{be}} - 0.69 \right) \quad (2-4)$$

where,

V_{be} = Effective binder content by volume, percent

V_a = In-situ air voids in the HMA mixture (%)

C_H = Thickness correction factor for bottom-up cracking as shown in Equation (2-5).

$$C_H = \frac{1}{0.000398 + \frac{0.003602}{1 + e^{(11.02 - 3.49H_{HMA})}}} \quad (2-5)$$

where,

H_{HMA} = AC layer thickness

Once the cumulative damage is calculated, the bottom-up fatigue cracking (%) can be estimated using the transfer function given in Equation (2-6).

$$FC_{Bottom} = \left(\frac{1}{60} \right) \left(\frac{C_4}{1 + e^{C_1 C_1^* + C_2 C_2^* \log(DI_{Bottom} \cdot 100)}} \right) \quad (2-6)$$

where,

FC_{Bottom} = Bottom-up fatigue cracking (in the percentage of area)

DI_{Bottom} = Cumulative damage at the bottom of the AC layer

C_1, C_2, C_4 = Transfer function coefficients where C_2 is a function of thickness for HMA thickness between 5 and 12 inches

C_1^* and C_2^* can be determined using Equation (2-7) and Equation (2-8).

$$C_1^* = -2C_2^* \quad (2-7)$$

$$C_2^* = -2.40874 - 39.748(1 + H_{HMA})^{-2.856} \quad (2-8)$$

Table 2-3 summarizes the local calibration coefficients for bottom-up cracking model among several states.

Table 2-3 Local calibration coefficients for bottom-up cracking

States	C_1	C_2	C_4	Standard deviation
Michigan	0.67	0.56	6000	$0.01 + \frac{32.913}{1 + e^{1.3972 - 0.9576 \times \log(D)}}$
Missouri	0.31	$C2 < 5'' = 1.367, C2 > 12'' = 2.067,$ $C2(5'' < \text{hac} < 12'') = 0.867 + 0.1 * \text{hac}$	6000	-
Georgia	2.2	2.2	6000	$1 + \frac{10}{1 + e^{7.5 - 6.5 \times \log(D + 0.0001)}}$
Louisiana	0.892	0.892	6000	-
Virginia	0.319	0.319	-	-
Tennessee	1.023	0.045	6000	-
Oklahoma (East Region)	3.26	-	6000	-
Oklahoma (West Region)	4.12	-	6000	-
Oklahoma (East region)	3.26	-	6000	-
Oklahoma (West region)	4.12	-	6000	-
Alabama	1	4.5	6000	$1.1 + \frac{22.9}{1 + e^{-0.1214 - 2.0565 \times \log(D + 0.0001)}}$
North Carolina	0.2437	0.24377	6000	-
Wyoming	0.4951	1.469	6000	-
Arkansas	0.688	0.294	6000	-
Colorado	0.07	2.35	6000	$0.01 + \frac{15}{1 + e^{-1.6673 - 2.4656 \times \log(D)}}$
New Mexico	0.625	0.25	6000	-
Oregon	0.560	0.225	6000	-
South Carolina	0.47	0.47	6000	-
Washington	1.071	1	6000	-
Pavement-ME v2.6	1.31	$C2 < 5'' = 2.1585,$ $C2 > 12'' = 3.9666,$ $C2(5'' < \text{hac} < 12'') = (0.867 + 0.25$ $83 * \text{hac}) * 1$	6000	$1.13 + \frac{13}{1 + e^{7.57 - 15.5 \times \log(D + 0.0001)}}$

2.3.1.2. Fatigue cracking (top-down)

Top-down or longitudinal cracking is a load-related distress due to repeated axle load. It appears in the form of cracks parallel to the wheel path and starts at the surface of the AC layer.

Old model: The damage calculation for top-down cracking is the same as bottom-up cracking for the old model except for the thickness correction factor and the transfer function, as shown in Equation (2-9) and Equation (2-10).

$$C_H = \frac{1}{0.01 + \frac{12.00}{1 + e^{(15.676 - 2.8186 H_{HMA})}}} \quad (2-9)$$

$$FC_{Top} = 10.56 \left(\frac{C_3}{1 + e^{C_1 - C_2 \log(DI_{Top})}} \right) \quad (2-10)$$

where,

FC_{Top} = Top-down fatigue cracking (in ft/mile)

DI_{Top} = Cumulative damage at the top of the AC layer

C_1, C_2, C_3 = Transfer function coefficients

New model: The new top-down cracking model is based on fracture mechanics concepts (48). It is expressed in percentage rather than ft./mile. The model involves crack initiation and propagation [based on Paris' law (49)]. Crack initiation is defined as a crack length of 7.5 mm (0.3 inches). Equation (2-11) shows the time to crack initiation formulated using regression over longitudinal and alligator cracking data from the LTPP database.

$$t_0 = \frac{K_{L1}}{1 + e^{K_{L2} \times 100 \times (a_0/2A_0) + K_{L3} \times HT + K_{L4} \times LT + K_{L5} \times \log_{10} AADTT}} \quad (2-11)$$

where,

t_0 = Time to crack initiation, days

H_T = Annual number of days above 32°C

L_T = Annual number of days below 0°C

$AADTT$ = Annual average daily truck traffic (initial year)

$a_0/2A_0$ = Energy parameter

$K_{L1}, K_{L2}, K_{L3}, K_{L4}, K_{L5}$ = Calibration coefficients for time to crack initiation

The top-down cracking is expressed in percentage using the transfer function, as shown in Equation (2-12).

$$L(t) = L_{MAX} e^{-\left(\frac{C_1 \rho}{t - C_3 t_0}\right)^{C_2 \beta}} \quad (2-12)$$

where,

$L(t)$ = Top-down cracking expressed as total lane area (%)

L_{MAX} = Maximum area of top-down cracking (%) – a value of 58% is assumed

t = Analysis month in days

ρ = Scale parameter for the top-down cracking curve as shown in Equation (2-13).

$$\rho = \alpha_1 + \alpha_2 \times \text{Month} \quad (2-13)$$

β = Shape parameter for the top-down cracking curve as shown in Equation (2-14).

$$\beta = 0.7319 \times (\log_{10} \text{ Month})^{-1.2801} \quad (2-14)$$

where,

α_1 and α_2 are functions of the climatic zone (wet freeze, wet non-freeze, dry freeze, dry non-freeze)

Table 2-4 summarizes the local calibration coefficients of the top-down cracking model. These coefficients have been obtained for the old top-down cracking model.

Table 2-4 Local calibration coefficients for top-down cracking

States	C_1	C_2	C_3	Standard deviation
Michigan	2.97	1.2	1000	$300 + \frac{3000}{1 + e^{7.5 - 6.5 \times \log(D_{bottom} + 0.0001)}}$
Tennessee	6.44	0.27	204.54	
Oklahoma (East Region)	6.6	4.6	723	-
Oklahoma (West Region)	6.1	4.23	723	-
Iowa	0.82	1.18	1000	-
Kansas	4.5	-	36000	-
Arkansas	3.016	0.216	1000	-
New Mexico	3	0.3	1000	-
Oregon	1.453	0.097	1000	-
South Carolina	0.2	0.1	3.97	-
Washington	6.42	3.596	1000	-
Pavement-ME v2.3	7	3.5	1000	-

2.3.1.3. Transverse (thermal) cracking model

Thermal cracking is associated with the contraction of the HMA material due to surface temperature fluctuations. The temperature variations affect the volume changes of the material. Consequently, stress develops due to the continual contraction of the materials and the restrained conditions, which causes thermal cracks. Typically, thermal cracking in flexible pavements occurs due to the temperature drop experienced by the pavement in cold conditions. A thermal crack will initiate when the tensile stresses in the HMA layers become equal to or greater than the material's tensile strength. The initial cracks propagate through the HMA layer with more thermal cycles. The amount of crack propagation induced by a given thermal cooling cycle is predicted using the Paris law of crack propagation. Experimental results indicate that reasonable estimates of A and n can be obtained from the indirect tensile creep-compliance and tensile strength of the HMA per Equations (2-15 and 2-16).

$$\Delta C = A(\Delta K)^n \quad (2-15)$$

where,

- ΔC = Change in the crack depth due to a cooling cycle
- ΔK = Change in the stress intensity factor due to a cooling cycle
- A, n = Fracture parameters for the HMA mixture

$$A = k_t \beta_t 10^{[4.389 - 2.52 \text{Log}(E_{HMA} \sigma_m \eta)]} \quad (2-16)$$

where,

- η = $0.8 \left[1 + \frac{1}{m} \right]$
- k_t = Regression coefficient determined through field calibration
- E_{HMA} = HMA indirect tensile modulus, psi
- σ_m = Mixture tensile strength, psi
- m = The m-value derived from the indirect tensile creep compliance curve measured in the laboratory
- β_t = Local or mixture calibration factor

The stress intensity factor, K , has been incorporated in the Pavement-ME through a simplified equation developed from theoretical finite element studies using the model shown in Equation (2-17).

$$K = \sigma_{tip} (0.45 + 1.99(C_o)^{0.56}) \quad (2-17)$$

where,

- σ_{tip} = Far-field stress from pavement response model at a depth of crack tip, psi
- C_o = Current crack length, feet

Equation (2-18) shows the transfer function for transverse cracking in the Pavement-ME.

$$TC = \beta_{t1} N(z) \left[\frac{1}{\sigma_d} \text{Log} \left(\frac{C_d}{H_{HMA}} \right) \right] \quad (2-18)$$

where,

- TC = Observed amount of thermal cracking, ft/500ft
- β_{t1} = Regression coefficient determined through global calibration (400)
- $N[z]$ = Standard normal distribution evaluated at $[z]$
- σ_d = Standard deviation of the log of the depth of cracks in the pavement (0.769), in.

C_d = Crack depth, in.

H_{HMA} = Thickness of HMA layers, in.

Table 2-5 summarizes the modified local calibration coefficients for the various states.

Table 2-5 Local calibration coefficients for the thermal cracking model

States	Level 1	Level 2	Level 3	Standard deviation
Michigan	0.75	-	4	Level 1 K: 0.4258*THERMAL +210.08 Level 3 K: 0.7737*THERMAL +622.92
Missouri	0.61	-	-	-
Oklahoma (East Region)	$3 \times 10^{-7} \times$ $MAAT^{4.0319} - 54$	-	-	-
Oklahoma (West Region)	$3 \times 10^{-7} \times$ $MAAT^{4.0319} - 23$	-	-	-
Arizona	1.5	0.5	1.5	Level 1 K: 0.1468*THERMAL +65.027 Level 2 K: 0.2841*THERMAL +55.462 Level 3 K: 0.3972*THERMAL +20.422
Colorado	7.5	-	-	Level 1 K: 0.1468*THERMAL +65.027
Minnesota	-	-	1.85	-
Montana	-	-	0.25	-
Pavement- ME v2.6	$3 \times 10^{-7} \times$ $MAAT^{4.0319}$	$3 \times 10^{-7} \times$ $MAAT^{4.0319}$	$3 \times 10^{-7} \times$ $MAAT^{4.0319}$	Level 1 K: 0.14*THERMAL +168 Level 2 K: 0.14*THERMAL +168 Level 3 K: 0.14*THERMAL +168

2.3.1.4. Rutting model

Due to axle loads, rutting is the total accumulated plastic strain in different pavement layers (AC, base/sub-base, and subgrade). It is calculated by summing up the plastic strains at the mid-depth of individual layers accumulated for each time increment. Equation (2-19) shows the permanent plastic strain for the AC layer.

$$\Delta_{p(HMA)} = \varepsilon_{p(HMA)} h_{HMA} = \beta_{1r} k_z \varepsilon_{r(HMA)} 10^{k_{1r}} T^{k_{2r}} \beta_{2r} N^{k_{3r}} \beta_{3r} \quad (2-19)$$

where,

$\Delta_{p(HMA)}$ = Permanent plastic deformation in the AC layer

$\varepsilon_{p(HMA)}$ = Accumulated permanent or plastic axial strain in the AC layer/sublayer

$\varepsilon_{r(HMA)}$ = Resilient or elastic strain calculated by the structural response model at the mid-depth of each AC sublayer

$h_{(HMA)}$ = Thickness of the AC layer/sublayer

N = Number of axle load repetitions

T = Pavement temperature

k_z = Depth confinement factor

k_{1r}, k_{2r}, k_{3r} = Global field calibration parameters

$\beta_{1r}, \beta_{2r}, \beta_{3r}$ = Local or mixture field calibration constants

The permanent plastic strain can be expressed for the unbound layers, as shown in Equation (2-20).

$$\Delta_{p(soil)} = \beta_{s1} k_{s1} \varepsilon_v h_{soil} \left(\frac{\varepsilon_o}{\varepsilon_r} \right) e^{-\left(\frac{\rho}{n} \right)^\beta} \quad (2-20)$$

where,

$\Delta_{p(Soil)}$ = Permanent plastic deformation for the unbound layer/sublayer

ε_o = Intercept determined from laboratory repeated load permanent deformation tests

n = Number of axle load applications

ε_r = Resilient strain imposed in laboratory tests to obtain material properties ε_o , β , and ρ

ε_v = Average vertical resilient or elastic strain in the layer/sublayer and calculated by the structural response model

h_{soil} = Unbound layer thickness

k_{s1} = Global calibration coefficients (different for granular and fine-grained material)

β_{s1} = Local calibration constant for rutting in the unbound layers (base or subgrade)

The total rutting is calculated based on Equation (2-21) below:

$$\text{Rut Depth}_{\text{Total}} = \Delta_{HMA} + \Delta_{\text{Base/subbase}} + \Delta_{\text{Subgrade}} \quad (2-21)$$

Table 2-6 presents the local calibration coefficients for different states.

2.3.1.5. IRI model (flexible pavements)

IRI is a measure of ride quality provided by a pavement surface and affects vehicle operation cost, safety, and driver comfort. The IRI model is based on findings from multiple studies showing that IRI at any age is a function of the initial construction ride quality and the development of different distresses over time that impact ride quality. IRI can be formulated using the initial IRI and distresses (fatigue cracking, transverse cracking, and rutting), as shown in Equation (2-22).

Table 2-6 Local calibration coefficients for the rutting model

States	β_{1r}	β_{2r}	β_{3r}	β_{qb}	β_{sq}	Standard deviation
Michigan	0.945	1.3	0.7	0.0985	0.0367	HMA: $0.1126 \cdot \text{RUT}^{0.2352}$ BASE: $0.1145 \cdot \text{RUT}^{0.3907}$ SG: $3.6118 \cdot \text{RUT}^{1.0951}$
Missouri	0.899	-	-	1.0798	0.9779	-
Georgia	-	-	-	0.5	0.3	HMA: $0.20 \cdot \text{RUT}^{0.55} + 0.001$
Louisiana	0.80	-	0.85	-	0.40	-
Virginia	0.664	-	-	0.151	0.151	-
Tennessee (Plain area)	0.111	-	-	0.196	0.722	-
Tennessee (Mountain area)	0.177	-	-	1.034	0.159	-
Oklahoma (East Region)	0.79	0.53	1.48	0.15	1.29	-
Oklahoma (West Region)	0.21	0.74	1.03	0.23	1.03	-
Arizona	0.69	1	1	0.14	0.37	HMA: $0.0999 \cdot \text{RUT}^{0.174} + 0.001$ BASE: $0.05 \cdot \text{RUT}^{0.115} + 0.001$ SG: $0.05 \cdot \text{RU}^{0.085} + 0.001$
Iowa	-	1.15	-	0.001	0.001	-
Kansas	0.9	-	-	-	0.3251	-
North Carolina	0.947	0.862	1.354	0.53767	1.5	-
Texas	2.39	-	0.856	-	0.5	-
Wyoming	-	-	-	0.4	0.4	-
Arkansas	1.20	1	0.8	1	0.5	-
Colorado	1.34	1	1	0.4	0.84	-
Montana	1.07	-	-	0.01	0.437	-
New Mexico	1.1	1.1	0.8	0.8	1.1	-
Ohio	0.51	-	-	0.32	0.33	-
Oregon	1.48	1.0	0.9	0	0	-
South Carolina	0.240	1	1	2.979	0.393	-
Utah	0.560	1	1	0.604	0.400	-
Washington	1.05	1.109	1.1	-	0	-
Pavement-ME v2.6	0.4	0.52	1.36	1	1	HMA: $0.24 \cdot \text{RUT}^{0.8026} + 0.001$ BASE: $0.1477 \cdot \text{RUT}^{0.6711} + 0.001$ SG: $0.1235 \cdot \text{RUT}^{0.5012} + 0.001$

$$IRI = IRI_o + C1(RD) + C2(FC_{Total}) + C3(TC) + C4(SF) \quad (2-22)$$

where,

IRI_o = Initial IRI at construction

FC_{Total} = Percent area of fatigue cracking (bottom-up), fatigue cracking (top-down), and reflection cracking in the wheel path

TC = Length of transverse cracking (including the reflection of transverse cracks in existing AC pavements)

RD = Average rut depth; $C1$, $C2$, $C3$, $C4$ = Calibration coefficients

SF = site factor, which can be expressed as shown in Equation (2-23) to Equation (2-25).

$$SF = (Frost + Swell) \times Age^{1.5} \quad (2-23)$$

$$Frost = \ln [(Rain + 1) \times (FI + 1) \times P_4] \quad (2-24)$$

$$Swell = \ln [(Rain + 1) \times (PI + 1) \times P_{200}] \quad (2-25)$$

where,

SF = Site factor

Age = Pavement age (years)

FI = Freezing index

PI = Subgrade soil plasticity index

$Rain$ = Mean annual rainfall

P_4 = Percent subgrade material passing No. 4 sieve

P_{200} = Percent subgrade material passing No. 200 sieve.

Table 2-7 presents the calibrated IRI coefficients in different states. Table 2-8 summarizes the distress thresholds for flexible pavements used in various states.

Table 2-7 Local calibration coefficients for the IRI model

States	C1	C2	C3	C4
Michigan	50.3720	0.4102	0.0066	0.0068
Missouri	58.9	0.3	0.0072	0.0129
Virginia	-	-	-	0.0392
Oklahoma (East Region)	5.23	0.127	0.013	0.0128
Oklahoma (West Region)	6.46	0.187	0.0098	0.023
Arizona	1.2281	0.1175	0.008	0.0280
Kansas	95	0.04	0.001	-
Colorado	35	0.3	0.02	0.019
New Mexico	-	-	-	0.015
Ohio	17.6	1.37	0.01	0.066
Pavement-ME v2.6	40	0.4	0.008	0.015

Table 2-8 Summary of design thresholds for flexible pavements

States	Bottom-up cracking (%)	Top-down cracking (ft/mile)	Total rutting	Thermal cracking	IRI
Michigan	20	-	0.5	1000	172
Missouri	10	-	0.50	1000	172
Louisiana	15	-	0.4	500	160
Virginia	10	-	0.4	500	160
Tennessee	10	2000	0.4	500	160
Oklahoma	20	-	0.4	630	169
Arizona	20	-	0.4	630	169
Kansas	20	-	0.4	630	169
Colorado	10	2000	0.4	1500	160

2.3.2 Performance Models for Rigid Pavements

2.3.2.1. Transverse cracking model

Transverse slab cracking in the Pavement-ME is calculated as the percentage of slabs cracked, including all severity levels. The mechanism involves independently predicting the bottom-up and top-down cracking and utilizing a probabilistic relationship to combine both, eliminating the possibility of both co-occurring. The fatigue damage for both bottom-up and top-down is defined using Miner's law as given in Equation (2-26):

$$DI_F = \sum \frac{n_{i,j,k,l,m,n,o}}{N_{i,j,k,l,m,n,o}} \quad (2-26)$$

where,

DI_F = Total fatigue damage (bottom-up or top-down)

$n_{i,j,k,l,m,n,o}$ = Actual load applications applied at age i , month j , axle type k , load level l , the equivalent temperature difference between top and bottom PCC surfaces m , traffic offset path n , and hourly truck traffic fraction o

$N_{i,j,k,l,m,n,o}$ = Allowable number of load applications applied at age i , month j , axle type k , load level l , the equivalent temperature difference between top and bottom PCC surfaces m , traffic offset path n , and hourly truck traffic fraction o

The allowable number of load applications is a function of PCC strength and applied stress and is calculated based on Equation (2-27):

$$\log(N_{i,j,k,l,m,n,o}) = C_1 \cdot \left(\frac{MR_i}{\sigma_{i,j,k,l,m,n,o}} \right)^{C_2} \quad (2-27)$$

where,

MR_i = Modulus of rupture of the PCC slab at the age i

$\sigma_{i,j,k,l,m,n}$ = Applied stress at the age i , month j , axle type k , load level l , the equivalent temperature difference between top and bottom PCC surface m , traffic offset path n , and hourly truck traffic fraction o

C_1, C_2 = Fatigue life calibration coefficients

The fraction of slabs cracked is predicted using Equation (2-28) for both bottom-up and top-down cracking:

$$CRK = \frac{1}{1 + C_4(DI_F)^{C_5}} \quad (2-28)$$

where,

CRK = Predicted fraction of bottom-up or top-down cracking

Once the bottom-up and top-down cracking is estimated, the percentage of slabs cracked is calculated using Equation (2-29).

$$TCRACK = (CRK_{\text{Bottom-up}} + CRK_{\text{Top-down}} - CRK_{\text{Bottom-up}} \cdot CRK_{\text{Top-down}}) \cdot 100 \quad (2-29)$$

where,

$TCRACK$ = Total transverse cracking (percentage of slabs cracked with all severities)

$CRK_{\text{Bottom-up}}$ = Predicted fraction of bottom-up transverse cracking

$CRK_{\text{Top-down}}$ = Predicted fraction of top-down transverse cracking

Table 2-9 summarizes the transverse cracking model local calibration coefficients in different states.

Table 2-9 Local calibration coefficients for the rigid transverse cracking model

States	C1	C2	C4	C5	Standard deviation
Michigan	-	-	0.23	-1.80	$1.34 \cdot CRK^{0.6593}$
Louisiana	2.75	-	1.16	-1.73	-
Idaho	2.366	1.22	0.52	-2.17	-
Arizona	-	-	0.19	-2.067	-
Minnesota	-	-	0.9	-2.64	-
South Carolina	1.25	1.22	-	-	-
Pavement-ME v2.6	2	1.22	0.52	-2.17	$3.5522 \cdot CRK^{0.3415} + 0.75$

2.3.2.2. Joint faulting model

The transverse joint faulting is calculated monthly in the Pavement-ME using the material properties, climatic conditions, present faulting level, pavement design properties, and axle loads

applied. Total faulting is the sum of faulting increments from previous months and is predicted using Equations (2-30) to (2-33) below.

$$\text{Fault}_m = \sum_{i=1}^m \Delta \text{Fault}_i \quad (2-30)$$

$$\Delta \text{Fault}_i = C_{34} \times (\text{FAULTMAX}_{i-1} - \text{Fault}_{i-1})^2 \times \text{DE}_i \quad (2-31)$$

$$\text{FAULTMAX}_i = \text{FAULTMAX}_0 + C_7 \times \sum_{j=1}^m \text{DE}_j \times \log(1 + C_5 \times 5.0^{\text{EROD}})^{C_6} \quad (2-32)$$

$$\text{FAULTMAX}_0 = C_{12} \times \delta_{\text{curling}} \times \left[\log(1 + C_5 \times 5.0^{\text{EROD}}) \times \log\left(\frac{P_{200} \times \text{WetDays}}{P_s}\right) \right]^{C_6} \quad (2-33)$$

where,

Fault_m = Mean joint faulting at the end of month m

ΔFault_i = Incremental change (monthly) in mean transverse joint faulting during the month i

FAULTMAX_i = Maximum mean transverse joint faulting for the month i

FAULTMAX_0 = Initial maximum mean transverse joint faulting

EROD = Erodibility factor for base/subbase

DE_i = Differential deformation energy of subgrade deformation accumulated during the month i

δ_{curling} = Maximum mean monthly slab corner upward PCC deflection due to temperature curling and moisture warping., P_s = Overburden pressure on the subgrade, P_{200} = Percent subgrade soil material passing No. 200 sieve

WetDays = Average annual number of wet days (greater than 0.1 in rainfall)

$C_{1,2,3,4,5,6,7,12,34}$ = Calibration coefficients

C_{12} and C_{34} are defined by Equation (2-34) and Equation (2-35):

$$C_{12} = C_1 + C_2 \times \text{FR}^{0.25} \quad (2-34)$$

$$C_{34} = C_3 + C_4 \times \text{FR}^{0.25} \quad (2-35)$$

FR = Base freezing index defined as the percentage of time (in hours) the top base temperature is below freezing (32 °F) temperature to the total number of hours in design life

Damage in a doweled joint for the current month is estimated using Equation (2-36).

$$\Delta \text{DOWDAM}_{\text{tot}} = \sum_{j=1}^N C_8 \times F_j \frac{n_j}{10^6 df_c^*} \quad (2-36)$$

where,

$\Delta DOWDAM_{tot}$ = Cumulative dowel damage for the current month

n_j = Number of axle load applications for the current increment and load group j for the current month

N = Number of load categories

f_c^* = Estimated PCC compressive stress

d = Dowel diameter

C_8 = Calibration constant

F_j = Effective dowel shear force induced by axle loading of load category j

The faulting model local calibration results for several states are summarized in Table 2-10.

Table 2-10 Local calibration coefficients for the faulting model

States	C1	C2	C3	C4	C5	C6	C7	C8	Standard deviation
Wyoming	0.5104	0.00838	0.00147	0.08345	5999	0.504	5.9293	-	$0.0831*FAULT^{0.3426}+0.00521$
Georgia	0.595	1.636	0.00217	0.00444	-	0.47	7.3	-	$0.07162*FAULT^{0.368}+0.00806$
Louisiana	1.5276	-	0.00262	-	-	0.55	-	-	-
Idaho	0.516	-	-	-	-	-	-	-	-
Arizona	0.0355	0.1147	0.00436	1.1E-07	20000	2.0389	0.1890	400	$0.037*FAULT^{0.6532}+0.001$
Kansas	-	-	0.00164	-	-	0.15	0.01	-	-
Michigan	0.4	-	-	-	-	-	-	-	$0.0442*FAULT^{0.2698}$
Wyoming	0.5104	0.00838	0.00147	0.08345	5999	0.504	5.9293	-	$0.0831*FAULT^{0.3426}+0.00521$
Pavement-ME v2.6	0.595	1.636	0.00217	0.00444	250	0.47	7.3	400	$0.07162*FAULT^{0.368}+0.00806$

2.3.2.3. IRI model (rigid pavements)

IRI in the Pavement-ME is a linear relationship between the IRI at construction and change in other distresses (transverse cracking, joint faulting, and joint spalling) over time. As a linear relationship of these factors, IRI can be expressed by Equation (2-37).

$$IRI = IRI_I + C1 \times CRK + C2 \times SPALL + C3 \times TFAULT + C4 \times SF \quad (2-37)$$

where,

IRI = Predicted IRI

IRI_I = Initial IRI at the time of construction

CRK = Percent slabs with transverse cracking (all severities).

$SPALL$ = Percentage of joints with spalling (medium and high severities).

$TFAULT$ = Total joint faulting cumulated per mi

C_1, C_2, C_3, C_4 = Calibration coefficients

SF = Site factor, which can be calculated as shown in Equation (2-38)

$$SF = AGE(1 + 0.5556 \times FI)(1 + P_{200}) \times 10^{-6} \quad (2-38)$$

where,

AGE = Pavement age

FI = Freezing index, °F-days.

P_{200} = Percent subgrade material passing No. 200 sieve.

The joint faulting and transverse cracking for IRI calculation are obtained using previously described models. The joint spalling is calculated as shown in Equation (2-39)

$$SPALL = \left[\frac{AGE}{AGE + 0.01} \right] \left[\frac{100}{1 + 1.005(-12 \times AGE + SCF)} \right] \quad (2-39)$$

where,

$SPALL$ = percentage joints spalled (medium- and high-severities)

AGE = pavement age since construction

SCF = scaling factor based on site-, design-, and climate-related variables, which is estimated as given in Equation (2-40)

$$SCF = -1400 + 350 \times ACPCC \times (0.5 + PREFORM) + 3.4f_c'^{0.4} - 0.2(FTcycles \times AGE) + 43h_{PCC} - 536WC_{PCC} \quad (2-40)$$

where,

$ACPCC$ = PCC air content

AGE = Time since construction

$PREFORM$ = 1 if preformed sealant is present; 0 if not

f_c' = PCC compressive strength

$FTcycles$ = Average annual number of freeze-thaw cycles

h_{PCC} = PCC slab thickness; WC_{PCC} = PCC water/cement ratio

The flexible pavement IRI local calibration coefficients for various states are summarized in Table 2-11. Table 2-12 shows threshold values used for different distresses in various states.

Table 2-11 Local calibration coefficients for rigid IRI model

States	C1	C2	C3	C4
Michigan	1.198	3.570	1.4929	25.24
Georgia	1.05	0.5417	1.85	33.8
Idaho	0.845	0.4417	1.4929	28.24
Virginia	9.55	172.55	-	-
Arizona	0.60	3.48	1.22	45.20
Iowa	0.04	0.04	0.07	1.17
Kansas	-	-	9.38	70
Ohio	0.820	3.7	1.711	5.703
Pavement-ME v2.6	0.8203	0.4417	1.4929	25.24

Table 2-12 Summary of design thresholds for rigid pavements

States	Transverse cracking (%)	Joint faulting (in)	IRI (in/mile)
Michigan	15	0.125	172
Missouri	-	-	172
Louisiana	10	0.15	160
Idaho	10	0.15	169
Virginia	10	0.15	160
Arizona	10	0.15	169
Kansas	10	0.15	169
Colorado	10	0.15	160
Minnesota	15	0.12	-

2.4 LOCAL CALIBRATION PROCESS

As mentioned, the Pavement-ME uses performance prediction models that are nationally calibrated based on pavement material properties, structure, climate, truck loading conditions, and data from the LTPP program (50). However, these models may not accurately predict pavement performance if the input properties and data used for calibration do not reflect the state's unique conditions. Therefore, it is recommended that each SHA evaluates how well the nationally calibrated models predict field performance. If the predictions are unsatisfactory, local calibration of the Pavement-ME models is recommended to improve the pavement performance predictions that reflect the state's specific field conditions and design practices. The local calibration process confirms that the prediction models can accurately predict pavement distress and smoothness and determines the standard error associated with the prediction equations. This section summarizes the local calibration process per the local calibration guide, 2010 (7) and MEPDG, 2015 (51).

Step 1: Selection of input levels

The hierarchical input level must be selected before local calibration. This depends on the availability of inputs in the local database and the agency's laboratory and field-testing capabilities. The selection of input levels is a critical step as it impacts the standard error of prediction.

Step 2: Develop an experimental plan and sampling strategy

The agency needs to develop a statistically sound and practical experimental plan and sampling template for this step. The sampling strategy should consider the local construction, design, and rehabilitation practices. The design matrix should include a wide range of traffic, materials, and climatic inputs.

Step 3: Assess the adequate sample size for each distress

A reasonable number of sections should be selected for calibration. The minimum sample size for any distress can be estimated using Equation (2-41).

$$n = \left(\frac{Z_{\alpha/2} \times \sigma}{e_t} \right)^2 \quad (2-41)$$

where,

$Z_{\alpha/2}$ = z-value from a standard normal distribution

n = Minimum number of pavement sections

σ = Performance threshold

e_t = Tolerable bias = $Z_{\alpha/2} \times SEE$

SEE = Standard error of the estimate

Step 4: Selection of pavement sections

This step involves selecting the pavement sections to populate the experimental matrix developed in Step 2. Selection should include local construction practices, sections with and without overlay, pavements with non-conventional materials, and replicates. To incorporate any time-dependent effects, a minimum of three measured distress data should be available over ten years. In case of section inadequacy, LTPP sections can be added to enhance the database.

Step 5: Get Pavement-ME inputs and measured distress data

The Pavement-ME inputs and the measured distress data must be extracted from the local agency database based on the hierarchical input level determined in Step 1. The performance data must be converted to the Pavement-ME compatible units if the agency measurements are different.

The average maximum distress from the selected sections should exceed 50% of the threshold design criteria to incorporate considerable distress in the calibration process. Any outliers in the performance data should be reviewed, considering the maintenance activities or changes in agency policies. Further field investigation can be conducted to resolve any discrepancies.

Step 6: Conduct field and forensic investigation

This step aims to collect any missing data and investigate any discrepancies in the input data available in the local database. The testing protocol to be followed should be in accordance with the agency's practices. At the end of this step, the agency should ensure that a reasonable number of samples remain in the experimental matrix.

Step 7: Validation of global model coefficients to local conditions

For this step, the global coefficients are used to predict each performance measure for all sections included in the experimental matrix. A reliability of 50% should be used for this step. The predicted values are compared with the measured ones to calculate the bias and SEE. A plot of predicted versus measured values is created for each distress to visualize the accuracy of predictions to a line of equality (LOE). For a good fit, the points should lie along the LOE. The measured distress $y_{Measured}$ and predicted distress $x_{Predicted}$ can be modeled as a linear model as shown in Equation (2-42) where m is the slope, and b_o is the intercept.

$$y_{Measured} = b_o + m \times x_{Predicted} \quad (2-42)$$

Three hypothesis tests are conducted to evaluate the reasonableness of the global model. If any of these hypotheses fail, the models are recalibrated for local conditions:

- There is no systematic bias between the measured and predicted distress [Equation (2-43)]. This can be tested using a paired t -test.

$$H_0: \sum(y_{Measured} - x_{Predicted}) = 0 \quad (2-43)$$

- The slope parameter m is 1, and the intercept parameter b_o is zero [Equations (2-44) and (2-45)].

$$H_0: m = 1.0 \quad (2-44)$$

$$H_0: b_o = 0 \quad (2-45)$$

Step 8: Eliminate the local bias for Pavement-ME models

This step should eliminate the local bias by systematically changing the model coefficients. The approach should be based on the overall bias, SEE between the predicted and measured values,

and the causes associated with them. The calibration coefficients should be incorporated into the calibration process if they depend on material property, site factor, or design features. Table 2-13 summarizes the calibration coefficients affecting the bias and standard error.

Table 2-13 Calibration coefficients eliminating standard error and bias (1)

Pavement Type	Distress	Eliminate Bias	Reduce Standard Error
Flexible	Total rut depth	$k_{1r}, \beta_{1r}, \beta_{s1}$	$k_{2r}, k_{3r}, \beta_{2r}, \beta_{3r}$
	Fatigue bottom-up cracking	k_1, C_2	k_2, k_3, C_1
	Fatigue top-down cracking	k_1, C_2	k_2, k_3, C_1
	Thermal cracking	β_{f3}, k_{f3}	β_{f3}, k_{f3}
	IRI	C_4	C_2, C_3, C_4
Rigid	Faulting	C_1	C_1
	Transverse cracking	C_1, C_4	C_2, C_5
	IRI - JPCP	J_4	J_1

Step 9: Estimate the standard error of the estimate

After the bias has been eliminated, the SEE is computed between the measured and predicted distress. This SEE must be compared with the global SEE. Table 2-14 shows the recommended value for SEE and bias for different models.

Table 2-14 Recommended values for tolerable bias and SEE (28)

Pavement Type	Distress/performance parameter	Bias	SEE
Flexible	Fatigue cracking (% total lane area)	1.5	5
	Rutting (inches)	0.075	0.2
	Thermal cracking (ft/mile) Thermal Reflection cracking	200	650
	IRI (inch/mile)	20	65
Rigid	Transverse cracking (% slabs cracked)	4	15
	Faulting (inch)	0.02	0.07
	IRI (inch/mile)	20	65

If the SEE is lower than recommended, the calibration coefficients can be accepted and used for design. The hypothesis tests given in step 7 must be validated before accepting the coefficients. If the SEE exceeds the global value, the agency can still accept the coefficients or move to step 10 to eliminate the standard error.

Step 10: Eliminate standard error of estimate (SEE)

If the standard error of the estimate calculated in step 9 is higher than the recommended global value, it should be eliminated in the local calibration process. The standard error should be estimated for each category of the experimental matrix to identify the effects of any input

parameter on the overall standard error. The coefficients resulting in the minimum standard error can be used for design purposes.

Step 11: Assessment of the calibration process

After the above ten steps have been performed to establish the local calibration coefficients, they should be examined for reasonableness within each category of the experimental matrix and at different reliability levels.

2.5 CONCEPT OF RELIABILITY

The Pavement-ME estimates the performance of a pavement using mechanistic models and transfer functions. Although these estimates are rational for pavement design purposes, the actual field measurements may show variability. This variability may come from the uncertainties in estimating the future traffic, material, and construction variability, measurement error, uncertainties due to the use of level 2 and 3 inputs, and errors associated with the model predictions. To incorporate all these variabilities, Pavement-ME uses a reliability-based design. Reliability for any prediction can be defined as the probability of getting a prediction lower than the threshold prediction over the design life, as shown in Equation (2-46).

$$\text{Reliability} = P[\text{distress at the end of design life} < \text{Critical distress}] \quad (2-46)$$

If 100 sections have been designed at 90% reliability, on average, ten of them may fail before the end of design life. Design reliability levels may vary by distress type and IRI or may remain constant for each. It is recommended that the same reliability be used for all performance indicators (51). Except for IRI, reliability for all other models is estimated using a relationship between the standard deviation of measured distress as the dependent variable and mean predicted distress as the independent variable. The basic assumption implies that the error in predicting the distress is normally distributed on the upper side of the prediction (not on the lower side or near zero values). Figure 2-2 shows an example of IRI prediction at 50% reliability (mean prediction), prediction at any desired reliability R , and are associated with the probability of failure. For 90 percent design reliability, the dashed curve at reliability R should not cross the IRI at the threshold criteria throughout the design analysis period. Failing to do so may lead to a failure at the required reliability and indicates that a design modification (such as a pavement thickness increase) should be applied.

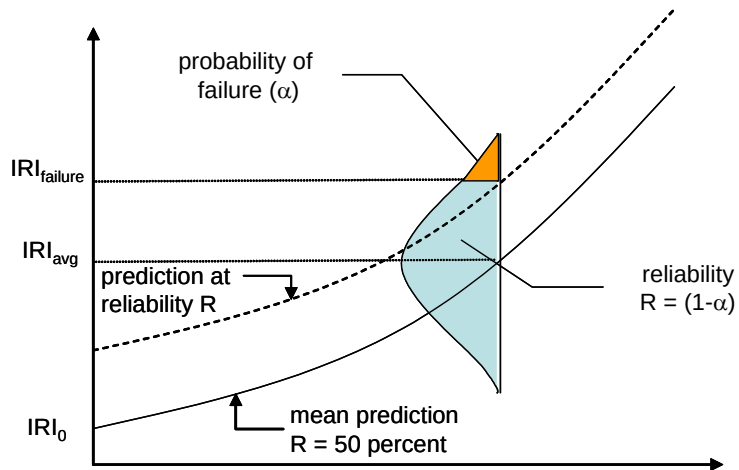


Figure 2-2 Design Reliability Concept for Smoothness (IRI)

2.6 IMPACT OF CALIBRATION ON PAVEMENT DESIGN

Several studies have been conducted to calibrate Pavement-ME transfer functions. Despite several calibration efforts, limited research is available on the effect of calibration on pavement design. Wu et al. (2014) calibrated the Pavement-ME models in Louisiana using Pavement-ME V1.3 (52). A total of 19 JPCP projects selected for this study had two base types: PCC over HMA and PCC over the unbound base. These 19 JPCP projects were designed using the Pavement-ME to estimate the effect on design thicknesses. The results showed that the Pavement-ME designs generated thinner PCC thicknesses (about 2 cm or 7%) compared to the AASHTO93 method (5). Tran et al. (2017) showed the effect of calibration on pavement design using the Missouri Department of Transportation (MoDOT) and Colorado Department of Transportation (CDOT) calibration results. One section, each for flexible and rigid pavement, was selected from existing MoDOT and CDOT projects. On average, the design thickness from local calibration was lower than that from the global model for both flexible and rigid sections (53). Mu et al. (2018) reviewed the effect of calibration on new JPCP design for seven states: Arizona, Colorado, Iowa, Louisiana, Missouri, Ohio, and Washington. The design thicknesses using global and local model coefficients were similar, such that five out of seven states had a difference of 13 mm or less. The Pavement-ME designs were thinner than AASHTO93 designs for high-traffic volume roads (by 50-70mm), whereas the thicknesses were similar for low-traffic volume roads. In rigid pavements, transverse cracking was the controlling distress for most cases

except for low-volume roads in Montana, where IRI was the critical distress (54). Singh et al. (2024) used the calibrated models in Michigan for pavement design to estimate the impact of calibration and for comparison with AASHTO93 designs. A total of 44 new flexible and rigid sections were designed. A comparison between AASHTO93 and Pavement-ME designs showed a reduction in HMA and PCC slab thicknesses for the latter approach. On average, the surface thicknesses using locally calibrated coefficients were thinner by 0.22 and 0.44 inches for flexible and rigid pavements, respectively. Critical design distresses for flexible pavements were bottom-up and thermal cracking. On the other hand, transverse cracking and IRI controlled the designs for rigid sections (55).

2.7 SENSITIVITY OF PAVEMENT-ME COEFFICIENTS

SHAs often struggle to identify the most critical data collection needs since the Pavement-ME requires several design inputs. Several studies have conducted sensitivity analyses to determine the most sensitive inputs to the distress prediction models for new and rehabilitated pavements to address this issue (56-62). However, limited research is available to determine the impact of each calibration coefficient on the predicted pavement distress and performance. Kim et al. (2014) conducted a sensitivity analysis for all the Pavement-ME models, determining the sensitivity by changing coefficients one at a time (26). This study performed the analyses using two in-service pavements representing typical Iowa's HMA and JPCP sections. Each calibration coefficient varied from its global value by 20% to 50%. For JPCP, the study concluded that the fatigue model-related calibration coefficients (C_1 and C_2) in the transverse cracking model are the most sensitive parameters. For the JPCP IRI model, coefficients C_1 (related to transverse cracking) and C_4 (related to site factor) are sensitive. Coefficient C_6 is the most sensitive for the faulting model. For flexible pavements, β_2 and β_3 are the most sensitive coefficients in fatigue cracking, whereas C_1 and C_2 are the most sensitive for IRI predictions. Dong et al. conducted a sensitivity analysis on calibration coefficients for the joint faulting model for JPCP sections in Ontario (12). The study also showed that C_6 is the most sensitive coefficient, followed by C_1 and C_2 . Both these studies quantified the sensitivity of coefficients using a sensitivity index (NSI) and a typical range of design inputs.

Parameter estimation is needed whenever a model is fitted to data to explain a phenomenon and is usually considered the same as curve-fitting or optimization. However, both

are distinctly different. While the optimization only focuses on minimizing the sum-of-squares or any other error criterion considering the parameters unimportant, parameter estimation also considers the parameters' errors (63). According to Beck and Arnold, parameter estimation is "a discipline that provides tools for the efficient use of data in the estimation of constants that appear in mathematical models and for aiding in modeling phenomena" (64).

Microsoft Excel's Solver® routine is used to estimate the parameters of a linear or nonlinear model but without computing the parameter errors, thus making it acceptable only for curve-fitting (63). However, according to Geeraerd et al., Solver® can accomplish parameter estimation if the sensitivity matrix is formulated and matrix multiplication is employed to compute the parameter errors (65). As per Dolan, the sensitivity matrix or Jacobian (J) is a matrix of the first derivatives of the model for each parameter and has the dimensions of n-by-p, where n and p are the numbers of data points and parameters, respectively (66). Thus, it is essential to know if any or all the parameters in a model are accurate and estimable, i.e., if they are statistically significant, they do not contain zero in the parameter confidence interval (CI). Hence, reporting the CI of any estimated parameter is equally important as the parameter errors.

Parameter identifiability depends on the scaled sensitivity coefficients (SSCs) and the minimization of the objective function (63). The SSCs can help determine whether a parameter is estimable and inform about its accuracy in terms of relative error. Several studies have used SSCs in various applications (other than pavements) to estimate the sensitivity of a parameter on a continuous scale of the independent variable (63, 66, 67).

The SSCs for the parameters are desired to be significant (the maximum value of SSC should be at least 10% of the largest value of the dependent variable) compared to the model η and uncorrelated with each other (63). The larger the SSC is for a parameter, the greater it will affect the model and the easier it will be to estimate. Moreover, the parameter with the largest SSC will also be the most accurate. However, suppose any of the SSCs are correlated, i.e., the ratio of SSCs of any two parameters is a constant (one is a linear function of the other); those parameters cannot be estimated together (only one can be calculated at a time) as the model η will respond to either of them identically. SSCs help assess a parameter's sensitivity on a continuous scale of the independent variable, highlight collinearity between coefficients, if any, and inform about the accuracy of the parameters, thus enhancing confidence in the parameter estimates. More importantly, determining SSC is a forward problem and does not require data, unlike NSI, which

requires the Pavement-ME design inputs (material, traffic, and climate). Overall, using SSC enhances confidence in the parameter estimates, leading to more reliable and informed decision-making in the analysis without data.

2.8 CHAPTER SUMMARY

This chapter summarized the calibration approaches, efforts, and transfer function coefficients from different states. Most states used the least squares method to calibrate the Pavement-ME coefficients. Least squares is a simplistic and popular approach based on the NIID assumption. These assumptions may not hold good for non-normally distributed data. Studies in different engineering fields have highlighted the advantages and applicability of the MLE method. This chapter also outlines the transfer functions for different flexible and rigid pavement models. A step-by-step approach for local calibration is described per the local calibration guide.

Pavement-ME uses a reliability-based design. The concept of reliability and its application in Pavement-ME design is explained. Several states have calibrated the Pavement-ME models to implement M-E design for local conditions. Despite several calibration efforts, the impact of calibration on pavement design has not been extensively evaluated. This chapter includes a literature review of studies that assessed the effect of calibration on pavement design. This consists of determining the design thicknesses and critical distress for pavement design. This chapter also includes a review of the sensitivity analysis of transfer function coefficients using the traditional NSI approach and describes the applicability of the SSC approach for sensitivity calculation. SSC has been widely used in different fields for parameter estimation and sensitivity calculations.

CHAPTER 3 - DATA FOR CALIBRATION

3.1 INTRODUCTION

This chapter discusses the inputs and performance data used for the local calibration process. A crucial step in local calibration involves choosing enough pavement sections that accurately represent the prevailing conditions in the area. The next step is to gather the necessary data for each of the selected pavement sections, including information on the pavement performance, maintenance history, and various Pavement-ME inputs (material, traffic, and climate) that directly influence performance predictions. The predictions are then compared to the actual performance of the constructed pavement sections. A pavement section refers to a specific stretch of road corresponding to a construction project, which may include up to two sections (such as different directions on a divided highway) with similar data inputs but varying measured pavement performance, traffic, and initial IRI. The accuracy of the predicted pavement performance in the Pavement-ME software depends on the information used to describe the in-service pavement. Thus, several inputs are essential for analyzing a particular pavement in the design software, particularly those with significant impacts on the expected performance. This chapter outlines the process for selecting pavement sections for local calibration and the steps in obtaining the required information for each pavement section.

First, the measured distresses from the MDOT PMS database were converted to Pavement-ME compatible units. Then, the time-series trends of all distress types were evaluated to identify potential projects for calibration. Also, these trends were explained, considering any significant maintenance activities over time. The information about maintenance activities over time will help to model a section in the Pavement-ME, i.e., whether an existing project should be considered a reconstruct or rehabilitated overlay project. The Pavement-ME inputs for these sections were also reviewed to obtain more updated or higher input levels. It's worth noting that a "project" refers to a specific job number in the construction records, while a "section" refers to multiple directions in a divided highway within a project. Hence, the number of sections is always greater than or equal to the number of projects. The project selection process, Pavement-ME inputs, and performance data have been summarized in this chapter.

3.2 MDOT PMS DATA

MDOT's Pavement Management System (PMS) and other available construction data sources were reviewed to identify the available input levels, units of measured performance data, and best possible estimates. The PMS and other sources were assessed to extract the following data:

- a. Performance data were evaluated for their measurement process and units and converted to the Pavement-ME compatible units (wherever required). Necessary assumptions were made for these conversions.
- b. The construction records, plans, job-mix formula (JMF), and other sources were used to identify the pavement cross-sections and material properties during construction. Any unavailable data was acquired from MDOT, or MDOT provided test results for the best possible estimates.
- c. Traffic data were collected from the construction records and MDOT Transportation Data Management System (TDMS). Level 2 data were used for traffic data based on road type, number of lanes, and vehicle class 9 traffic percentage.
- d. For Asphalt concrete (AC) mix and binder properties, DYNAMOD software was used, which is based on laboratory tests for Michigan mixes. The most common construction materials in Michigan were used for base, subbase, and subgrade properties.
- e. For climatic data, the updated NARR files for Michigan have been used (68).

3.2.1 Pavement Condition Measures Compatibilities

MDOT provided the PMS data from 1992 to 2019 (sensor data from 1998 to 2019). Biannually, MDOT obtains performance data on their pavement network by utilizing distress and laser-based measurements (sensors) for a 0.1-mile section. The information gathered on pavement distress in MDOT's PMS is categorized by distinct principle distress (PD) codes, where each PD code corresponds to a specific distress type (69). This pavement performance data was extracted for the selected projects and converted to Pavement-ME compatible units (where needed). In addition, MDOT personnel explained the distress calls made for the 2012 – 2017 data were only at the sampled locations (about 29.41% of any 0.1-mile segment of each control section). Therefore, it was suggested that a 0.2941 division factor be considered for those years of measured PMS data.

3.2.1.1. Selected distresses

The MDOT PMS and sensor database were carefully analyzed, and relevant data were extracted to obtain the required distress information. The current distress manual of MDOT PMS was used to determine all the principle distress (PD) codes corresponding to the predicted distresses in the Pavement-ME. The earlier versions of the PMS manual were also reviewed to ensure accurate data was extracted for all the years. The necessary steps for PMS data extraction include:

1. Identify the PDs that correspond to the Pavement-ME predicted distresses
2. Extract PDs and sensor data for each project
3. Convert (if necessary) MDOT PDs to the units compatible with the Pavement-ME
4. Summarize time-series data for each project and each distress type

Tables 3-1 and 3-2 summarize the identified and extracted pavement distresses and conditions for flexible and rigid pavements. This section also presents a detailed discussion of the conversion process for both flexible and rigid pavements.

Table 3-1 Flexible pavement distress measurement by MDOT

Flexible pavement distress	MDOT principle distresses (PDs)	MDOT units	Pavement-ME units	Conversion needed?
IRI	Directly measured	in/mile	in/mile	No
Top-down cracking	204, 205, 724, 725 , 501	miles	% area	Yes
Bottom-up cracking	234, 235, 220, 221 , 730, 731 , 501	miles	% area	Yes
Thermal cracking	101 , 103, 104, 114, 701, 703, 704 , 110, 501	No. of occurrences	ft/mile	Yes
Rutting	Directly measured	in	in	No
Reflective cracking	No specific PD	None	% area	N/A

Note: Bold numbers represent older PDs that are not currently in use; PD code 501 = No distress

Table 3-2 Rigid pavement distress measurement by MDOT

Rigid pavement distresses	MDOT principle distresses	MDOT units	Pavement-ME units	Conversion needed?
IRI	Directly measured	in/mile	in/mile	No
Faulting	Directly measured	in	in	Yes
Transverse cracking	112, 113, 501	No. of occurrences	% slabs cracked	Yes

Note: PD code 501 = No distress

3.2.1.2. Pavement distress unit conversion for HMA designs

It should be noted that the Pavement-ME predicted distresses for the local calibration were only considered. The corresponding MDOT PDs were determined and compared with distress types predicted by the Pavement-ME to verify if any conversions were necessary. MDOT measures pavement distresses related to HMA pavements are listed in Table 3-1. PD code 501 corresponds to no distress condition and has been used in all distresses except rutting and IRI. The conversion process (if necessary) for all distress types is as follows:

IRI: The IRI measurements in the MDOT sensor database are compatible with those in the Pavement-ME. Therefore, no conversion or adjustments were needed, and data could be used directly.

Top-down cracking: Top-down cracking is load-related longitudinal cracking in the wheel path. The PDs 204, 205, 724, and 725 were assumed to correspond to the top-down cracking in the MDOT PMS database because those may not have developed an interconnected pattern that indicates alligator cracking. Those cracks may show an early stage of fatigue cracking, which could also be bottom-up. Since estimating such cracking based on the PMS data is difficult, these cracks were converted to % area crack and then categorized into bottom-up or top-down cracking based on the thicknesses. The PDs are recorded in miles and need conversion to % area. Data from the wheel paths were summed into one value and divided by the total project length, as shown in Equation (3-1). The lane width was assumed to be 12 ft. The typical wheel path width of 3 feet was assumed as recommended by the LTPP distress identification manual (70).

$$\% AC_{top-down} = \frac{\text{Length of cracking (miles)} \times \text{width of wheelpaths (feet)}}{\text{Length of section (miles)} \times \text{Lane width (feet)}} \times 100 \quad (3-1)$$

Literature shows that the AC thickness determines whether the crack initiates from the bottom or the top. Therefore, top-down cracking can be a primary distress based on AC layer thickness. The calculated top-down cracking using Equation (3-1) is assigned as either bottom-up or top-down based on the total AC layer thickness. If the thickness exceeds a certain threshold, the cracking is considered top-down cracking; otherwise, it is categorized as bottom-up cracking. These thicknesses were obtained by a mechanistic approach using Mechanistic Empirical Asphalt Pavement Analysis (MEAPA) software. MEAPA was run for different surface types using typical MDOT design inputs, and damage was calculated for the first 12 months for a

single axle load of 9000 lb. Threshold thicknesses were determined where the tensile strain at the top of the AC layer is higher than at the bottom. Table 3-3 presents the minimum threshold thicknesses for top-down cracking for each fix type.

Table 3-3 Minimum thicknesses for top-down cracking

Fix type	Threshold thickness (in)
HMA overlay on rubblized concrete	6
HMA overlay on crushed and shaped HMA	4
New or reconstruct	5

Bottom-up cracking: Bottom-up cracking is alligator cracking in the wheel path. The PDs 234, 235, 220, 221, 730, and 731 match this requirement in the MDOT PMS database. The PDs have units of miles; however, to make those compatible with the Pavement-ME alligator cracking units, conversion to the percent of the total area is needed. This can be achieved by using the following Equation (3-2):

$$\%AC_{bottom-up} = \frac{\text{Length of cracking (miles)} \times \text{width of wheelpaths (feet)}}{\text{Length of section (miles)} \times \text{Lane width (feet)}} \times 100 \quad (3-2)$$

The widths of each wheel path and lane were assumed to be 3 feet and 12 feet, respectively. The LTPP distress identification manual recommends a typical wheel path width of 3 feet (70).

Thermal cracking: Thermal cracking corresponds to transverse cracking in flexible pavements. The transverse cracking is recorded as the number of occurrences, but the Pavement-ME predicts thermal cracking in feet/mile. To convert transverse cracking into feet/mile, the number of occurrences was multiplied by 3 feet for PDs 114 and 701 because these PDs are defined as "tears" (short cracks) that are less than half the lane width. For all other PDs, the number of occurrences was multiplied by the lane width (12 ft). All transverse crack lengths were summed and divided by the project length to get feet/mile, as shown in Equation (3-3).

$$TC = \frac{\sum \text{No. of Occurrences} \times \text{Lane Width (ft)}}{\text{Section Length (miles)}} \quad (3-3)$$

Thermal cracking predictions in the Pavement-ME are restricted to a maximum value of 2112 ft/mile due to a minimum crack spacing limit of 30 feet. This means Pavement-ME predictions at 50% reliability cannot exceed 2112 ft/mile. Due to this limitation and ARA recommendations, a 2112 ft/mile cutoff was decided where any measured data for a section above 2112 ft/mile was not used for calibration.

Rutting: This is the total amount of surface rutting all the pavement layers and unbound sub-layers contribute. The average rutting (left & right wheel paths) was determined for the entire project length. No conversion was necessary. It is assumed that the measured rutting corresponds to the total surface rutting predicted by the Pavement-ME.

3.2.1.3. Pavement distress unit conversion for rigid designs

For rigid sections, transverse cracking requires unit conversion. For all other distresses, MDOT records them in the Pavement-ME compatible units. Table 3-2 summarizes the distresses related to rigid sections, and the conversion process is discussed below:

IRI: The IRI in the MDOT sensor database does not need any conversion; the values were used directly.

Faulting: In the Pavement-ME, faulting is predicted as average per joint. MDOT's sensor data records the number of faults (FaultNum), average faulting (avgFault), and the maximum faulting (FaultMax) for every 0.1-mile segment. The faulting values had some inconsistencies. For the years between 2000 and 2011, faulting values are maximum fault callouts only (not average values). For 2012 and after, both average and maximum fault values are available. A correlation was developed between the maximum and average faulting values using data from 2013 to 2017 to resolve this issue. These correlations were used to estimate the average faulting from 2000 to 2011. Table 3-4 shows the regression equations between average and maximum faulting using the data from 2013 to 2017. These equations are based on the number of faults. It is important to note that ideally, the number of faults cannot be greater than the number of joints, but the number of faults in the database has records where they are more than the number of joints. These pseudo-fault values might come from cracking, spalling, bridge segments, etc. Therefore, the maximum number of fault counts was restricted to 36, and the average faulting to 0.4 inches to address this issue. Accordingly, any 0.1-mile section above these restricted faulting values was omitted from the calibration data.

Table 3-4 Correlation equations based on the number of faults

FaultNum		Equation (y is avgFault, x is FaultMax)	R-squared (2013-2017data)
From	To		
0	1	y=x	1
2	4	y = 0.3438x + 0.03	0.7189
5	40	y = 0.2132x + 0.0377	0.6074
41	ALL	y = 0.0936x + 0.0777	0.2476

The average joint faulting is calculated based on the number of faulting in a 0.1-mile section. It is assumed that if the number of faults is less or equal to the number of joints, faulting occurs at the joints only. In that case, the faulting unit conversion equation is as shown in Equation (3-4). If, for any 0.1-mile section, the number of faults is greater than the number of joints, that section is removed (cut) from the calibration data, as previously mentioned.

$$Fault = \frac{FAULnum \times FAULi}{N_{joints}} \quad (3-4)$$

where,

FAULnum = Number of faults in a 0.1 mile

FAULi = (FAULT_(Avg_Right) + FAULT_(Avg_Left))/2 = Average faulting in a 0.1 mile (inches)

N_{joints} is the number of joints in 0.1-mile (528 ft) segments, i.e., N_{joints}=528/Joint Spacing.

Transverse cracking: The transverse cracking distress is predicted as the percentage of slabs cracked in the Pavement-ME. However, MDOT measures transverse cracking as the number of transverse cracks. PDs 112 and 113 correspond to transverse cracking. The estimated transverse cracking must be converted to the percent slabs cracked using Equation (3-5).

$$\% \text{ Slabs Cracked} = \frac{\sum PD_{112,113}}{\left(\frac{\text{Section Length (miles)} \times 5280 \text{ ft}}{\text{Joint Spacing (ft)}} \right)} \times 100 \quad (3-5)$$

3.2.2 Condition Database for Local Calibration

Customized databases were created to efficiently analyze the condition of selected Pavement Distresses (PDs), which included distress and sensor data for multiple years. These databases were compiled using Microsoft Access and allowed for easy extraction of relevant data for projects of any length. The PMS condition data from 1992 to 2019 and sensor data from 1998 to 2019 were included in these databases. MATLAB codes were used to extract performance data

for a section of the given length. For divided highways, which can have an increasing and decreasing direction to indicate north/south or east/west bounds, both directions were included in the time-series data and considered separate sections. In contrast, distress data was collected in one direction for undivided highways.

3.3 PROJECT SELECTION CRITERIA

For local calibration, selecting in-service pavement sections that represent local pavement design, currently used materials, construction practices, and performance is essential. A set of project selection criteria was established to identify and choose these representative pavement sections. This approach ensured that the selected pavement sections met the required standards and could accurately represent Michigan's pavement network. The process for identifying and selecting pavement sections consists of the following steps:

1. Determine the minimum number of pavement sections required for calibration based on the statistical requirements.
2. Identify all available in-service pavement projects.
3. Extract all pavement distresses (pavement condition data) from the customized database for all identified projects in Step 2.
4. Evaluate the measured performance for all the identified projects.
5. Identify projects with adequate data, age, trend, and the Pavement-ME inputs available to develop a refined list.

3.3.1 Identify the Minimum Number of Required Pavement Sections

The MEPDG local calibration guide provides a method to evaluate the minimum number of required sections for each distress type. The minimum number of sections was calculated using Equation (3-6), and the results are summarized in Table 3-5 for each condition measure. The total number of projects available in Table 3-5 are combined projects from the previous calibration study (10) and newly selected projects from the current calibration effort.

$$n = \left(\frac{Z_{\alpha/2} \times \sigma}{e_t} \right)^2 \quad (3-6)$$

where;

$Z_{\alpha/2}$ = The z-value from a standard normal distribution

n = Minimum number of pavement sections

σ = Performance threshold

e_t = Tolerable bias $Z_{\alpha/2} \times SEE$

SEE = Standard error of the estimate

Table 3-5 Minimum number of sections for local calibration

Performance Model	Nationally calibrated SEE	Z ₉₀	Threshold	N (required number of sections)	Number of sections used	Total number of projects available
Flexible Pavements						
Fatigue, bottom-up (%)	5.01	1.64	20%	16	78	163
Fatigue, top-down (ft/mile or %)	583		2000 or 20%	12	133	
Thermal cracking (ft/mile) ¹	-		1000	-	133	
Rutting (in)	0.107		0.5	22	200	
IRI (in/mile)	18.9		172	83	178	
Rigid Pavements						
Transverse cracking (%)	4.52	1.64	15	11	48	46
Joint faulting (in)	0.033		0.125	14	79	
IRI (in/mile)	22		172	61	48	

Note: Fatigue top-down has been updated in the recent Pavement-ME V2.6. It is expressed in ft/mile for the old model and in % for the updated model.

N = minimum number of samples required for a 90% confidence level

1. No SEE , threshold, or N was reported for thermal cracking in the literature

3.3.1 Initial Projects Selection

The common pavement types in Michigan include:

1. HMA reconstruct
2. HMA over crush & shaped existing HMA
3. HMA over rubblized PCC
4. JPCP reconstruct

It is important to note that HMA over crushed and shaped existing HMA and HMA over rubblized existing PCC projects were analyzed as new reconstructed pavement. Sections were selected for the local calibration based on performance trends and to accommodate a wide range of different inputs, including layer thicknesses, traffic, region, etc.

MDOT provided a comprehensive database consisting of all the projects constructed in Michigan. Initially, all existing projects used in previous calibration efforts were reviewed, and additional performance data were extracted where possible. Additional projects were identified that can be potential candidates for the current local calibration effort. The PMS data extraction was completed for all required distress types in a compatible format with the Pavement-ME software. The time series for each pavement section's performance measures was observed to finalize the preliminary list of new potential candidate projects. To ensure a robust and appropriate set of data, the criteria used to identify additional performance data and the selection of new potential pavement projects include:

- The pavement section must have at least three measured data points over time. There are some exceptions to this criterion. Bottom-up cracking has relatively fewer data points; some sections with even two points have been included, considering further data points will be collected in the future. The same process was followed for transverse cracking in rigid sections. As previously explained, joint faulting and thermal cracking have been cut at specific values, so these data points are omitted from the calibration database.
- At least one of the distresses should have an increasing trend. Any section with decreasing and no or flat trends over time was excluded from the list.
- The previous maintenance history was observed for all pavement sections to explain any decrease or flat trend in the time series plot. If there were any major rehabilitation or reconstruction activities, the measured data from the year traffic opened initially to the very last year until the major repair took place are considered.
- The last recorded point should have a Distress Index (DI) of at least 5 for a section. DI is calculated by taking a weighted average of different distress types. DI was observed and limited to ensure sufficient distress for calibration and to capture adequate pavement performance trends.

Figures 3-1 and 3-2 illustrate examples of distress progressions for a selected and omitted flexible pavement section. The top-down cracking for the initial project selection was evaluated in feet/mile and later converted to a percentage. Similarly, Figures 3-3 and 3-4 present examples of the selected and omitted rigid pavement sections. The vertical dashed red line is the last reported construction, whereas the dotted blue line in the DI plot indicates reported maintenance activities. For example, Figure 3-1 shows the vertical dotted blue line in the DI plot that shows a

cold mill and resurface (CM&R) treatment was applied in 2012. In the same figure, the effect of this rehabilitation event can be noticed with a drop in measured distress in individual distress plots. Therefore, in this case, pavement section performance can be considered from 2001 to 2011. It should be noted that generally, minor maintenance [e.g., crack treatment (CT) or joint sealing (JS)] does not affect the time series trend since these minor maintenances represent non-structural fixes. Note that time series plots for rutting show a consistent drop in the 2012-2013 collection years, regardless of whether any maintenance is reported. This is likely due to changes in the data collection process or vendor differences.

Based on the criteria mentioned above, a total of 256 flexible sections and 88 rigid sections were initially selected. The performance of the chosen pavement sections was compared with all sections available in the MDOT database (2081 flexible sections and 442 rigid sections) to verify if the chosen sections represent the overall pavements in Michigan. Sections with at least three available data points are considered. Each section was categorized as good, fair, or poor performing based on the performance trend lines modified to reflect Michigan conditions (10). These trend lines are available only for bottom-up cracking, total rutting, and IRI for flexible sections, as well as transverse cracking and IRI for rigid sections. The performance categories depend on the measured performance trend relative to the reference lines. If the measured performance is below the good performance line, it is categorized as a good performing section, between the good and poor line, as fair, and above the poor performance line, as the poor performing section. The performance category was decided based on a previous calibration study (10). When the performance trend passes through more than one category zone, the zone with the maximum points is considered the performance category for that section. Also, the low-performance category is selected in case of an equal number of points for two different categories. Figures 3-5 and 3-6 show example sections for good, fair, and poor categories for IRI performance for flexible and rigid sections, respectively.

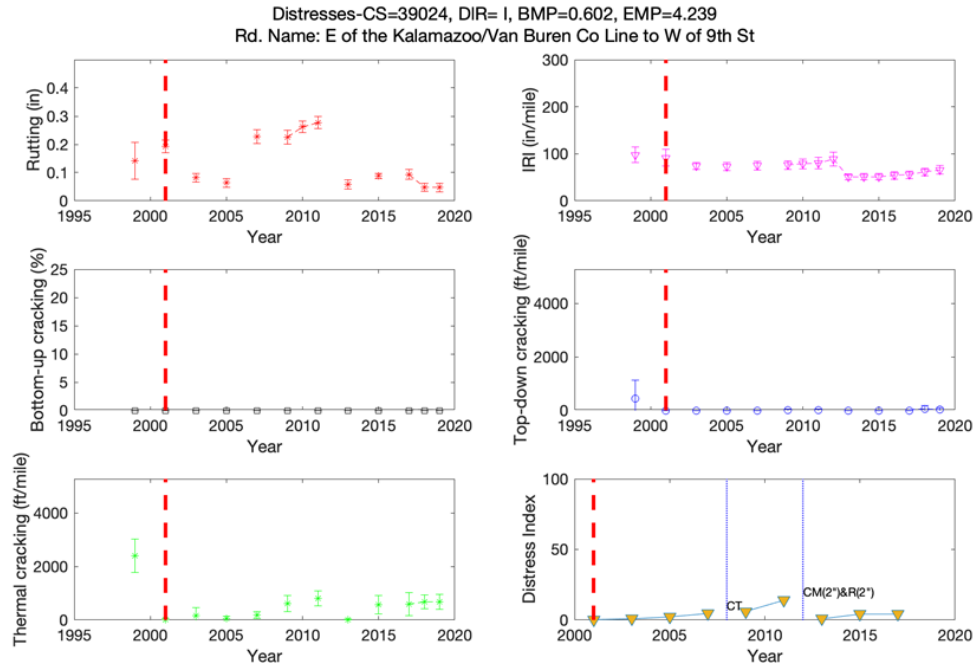


Figure 3-1 Example of selected flexible section

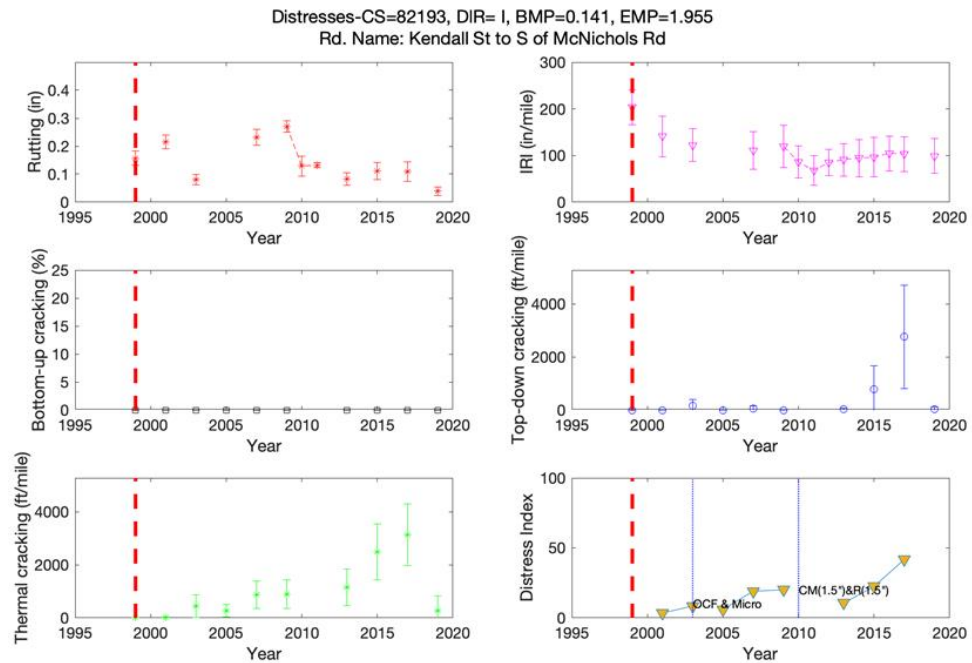


Figure 3-2 Example of an omitted flexible section

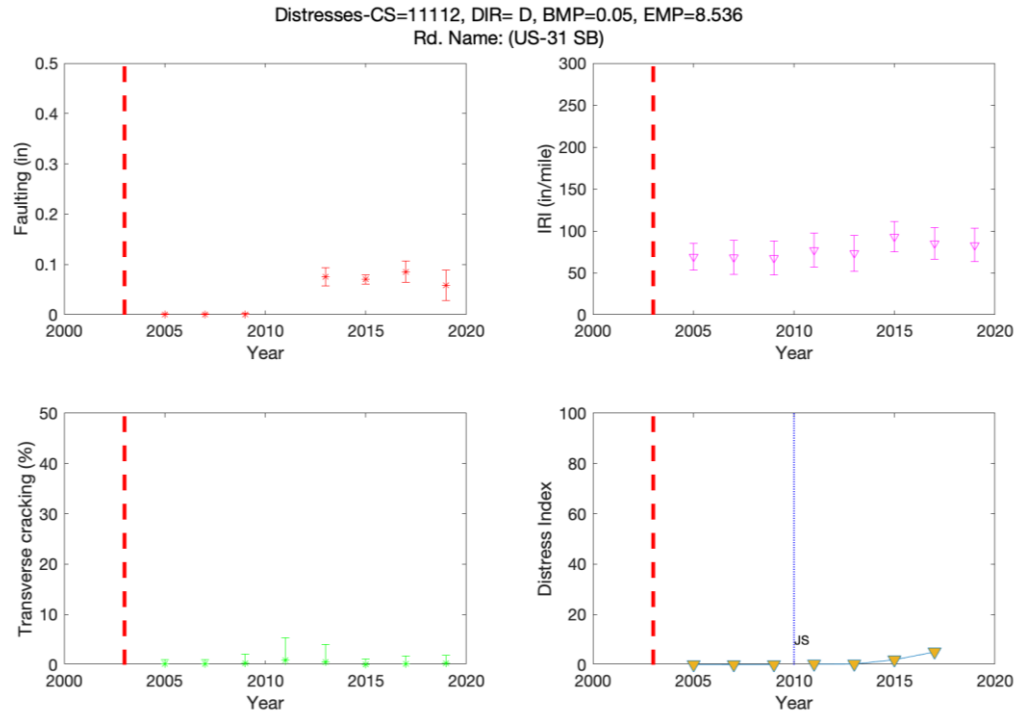


Figure 3-3 Example of a selected rigid section

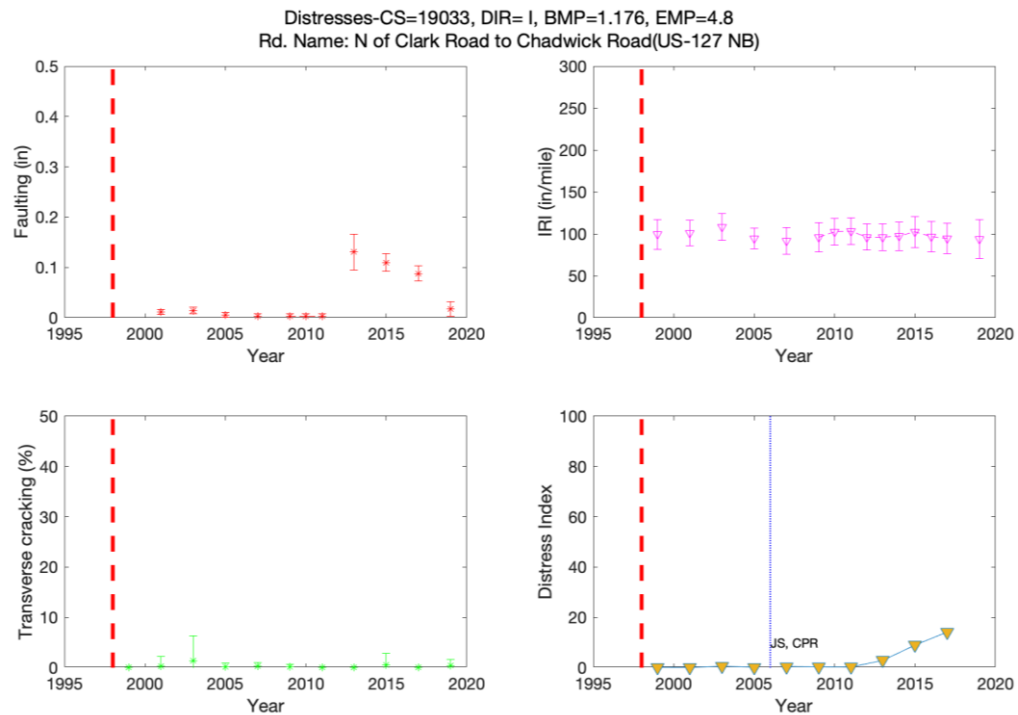


Figure 3-4 Example of an omitted rigid section

A similar method was followed for categorizing sections based on all other distresses. Figures 3-7 and 3-8 show the distribution of good, fair, and poor sections for rigid and flexible sections based on different distress criteria. Figures 3-7 and 3-8 show that the selected sections satisfactorily represent MDOT all sections for both flexible and rigid pavements.

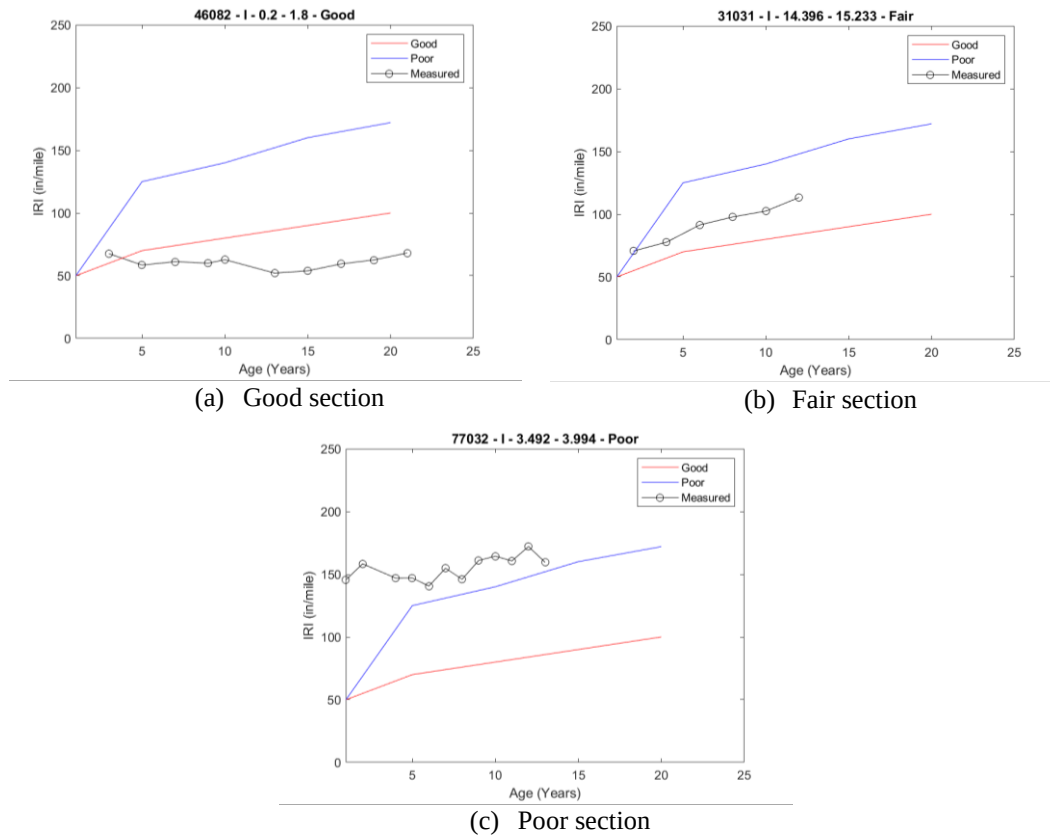


Figure 3-5 Categorization of flexible sections based on performance trends

The initially selected projects were further refined based on performance, availability of inputs, and initial IRI. The performance data for these initially selected sections is the average for the entire section length. This data is calculated by averaging the performance for every 0.1-mile segment in the project length. Data for every 0.1 mile has been reviewed to estimate performance data extent and reasonableness. Figures 3-9 to 3-13 show performance data for every 0.1-mile segment with years for all flexible sections.

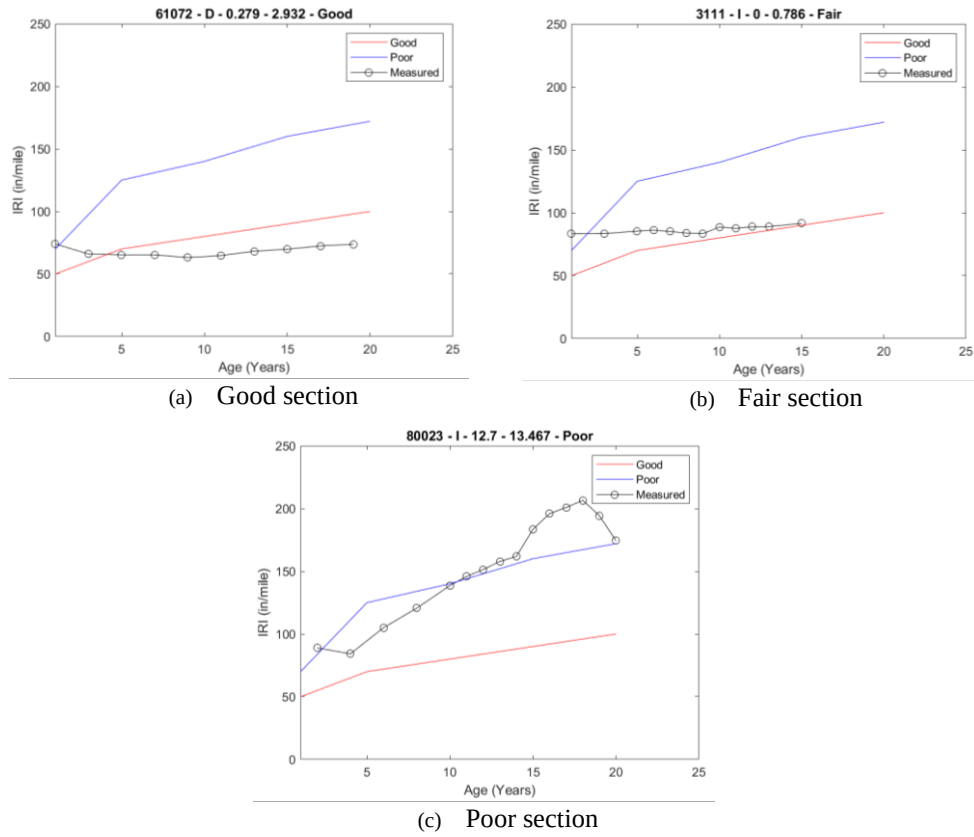
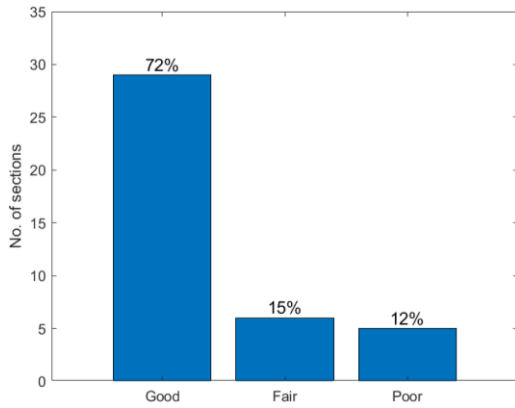
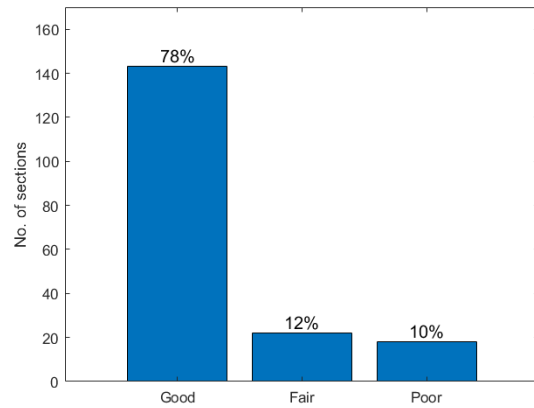


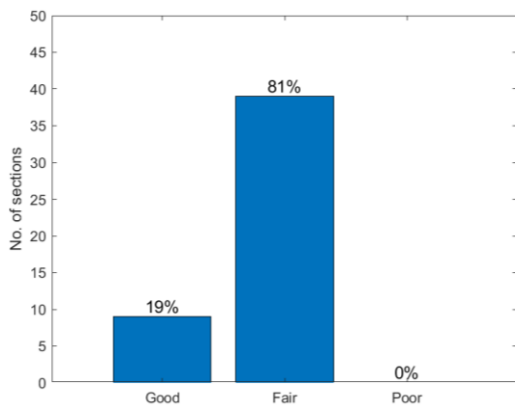
Figure 3-6 Categorization of rigid sections based on performance trends



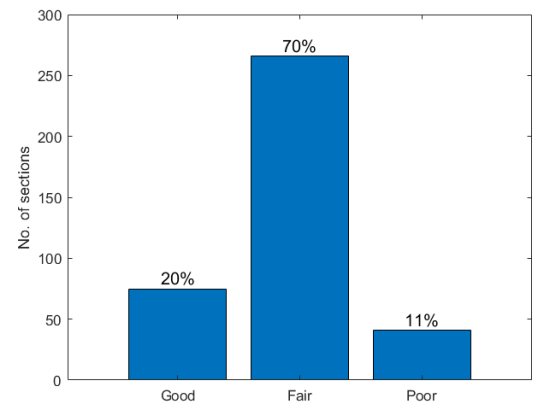
(a) Transverse cracking (selected sections)



(b) Transverse cracking (MDOT all sections)

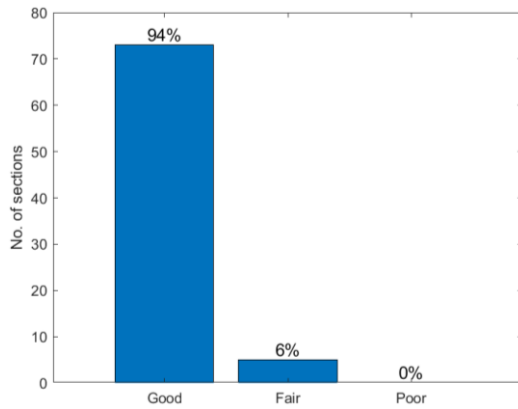


(c) IRI (selected sections)

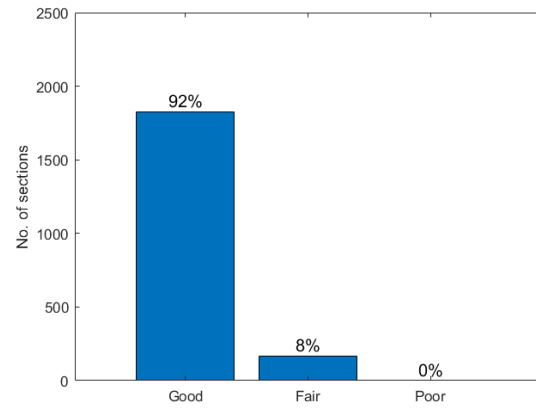


(d) IRI (MDOT all sections)

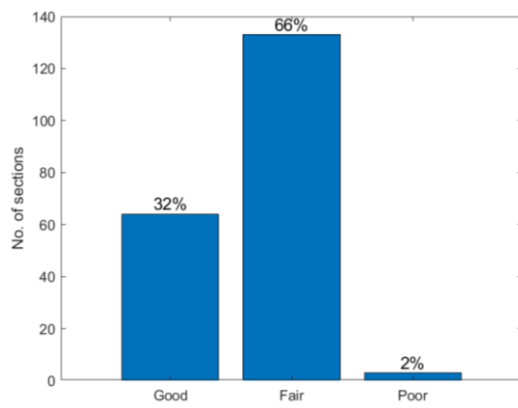
Figure 3-7 Comparison of selected rigid sections with all MDOT sections



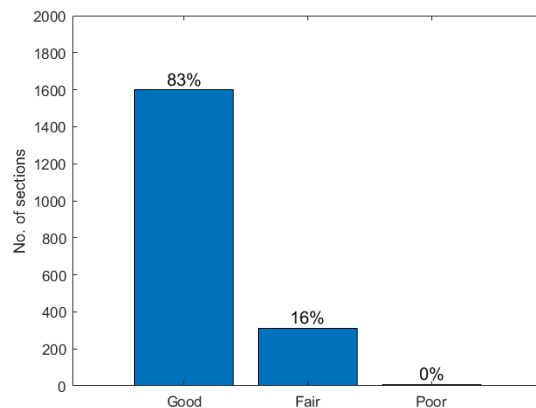
(a) Bottom-up cracking (selected sections)



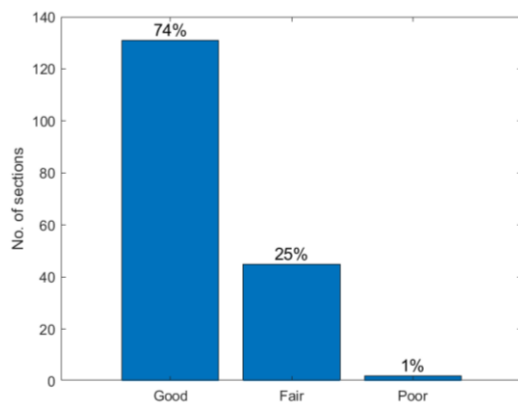
(b) Bottom-up cracking (All MDOT sections)



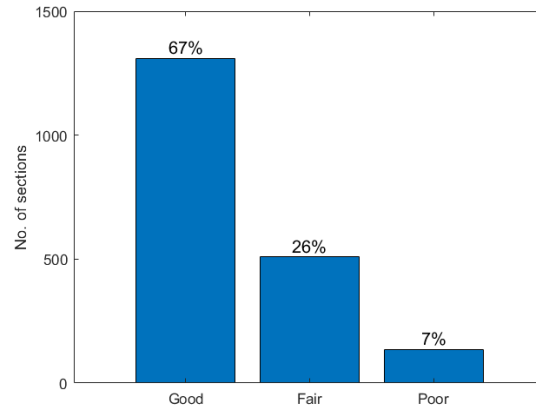
(c) Total rutting (selected sections)



(d) Total rutting (All MDOT sections)



(e) IRI (selected sections)



(f) IRI (All MDOT sections)

Figure 3-8 Comparison of selected flexible sections with all MDOT sections

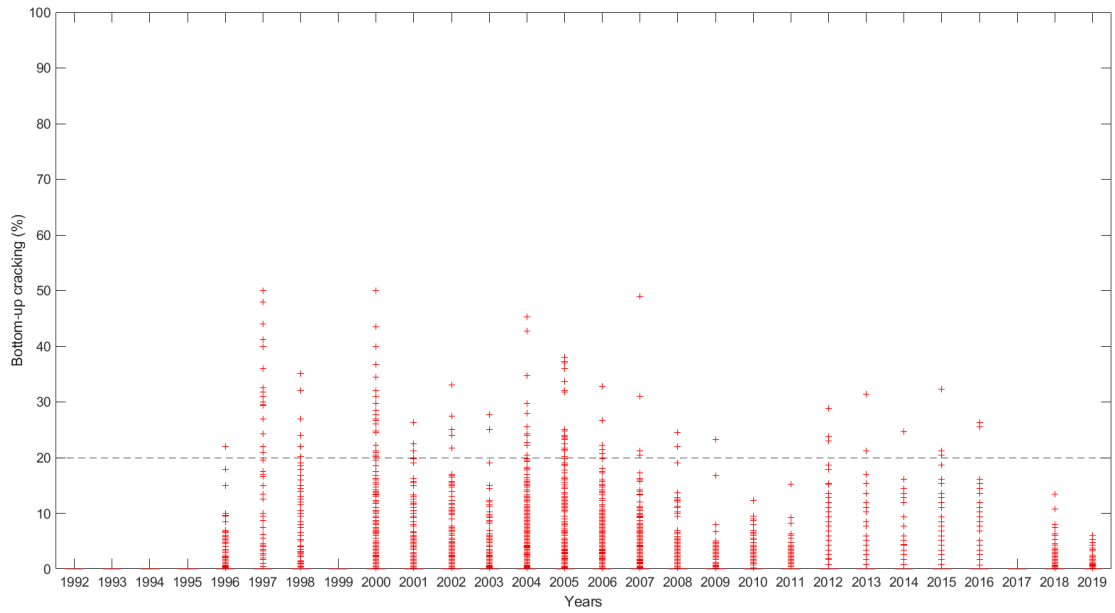


Figure 3-9 Bottom-up cracking at every 0.1-mile segment for flexible sections

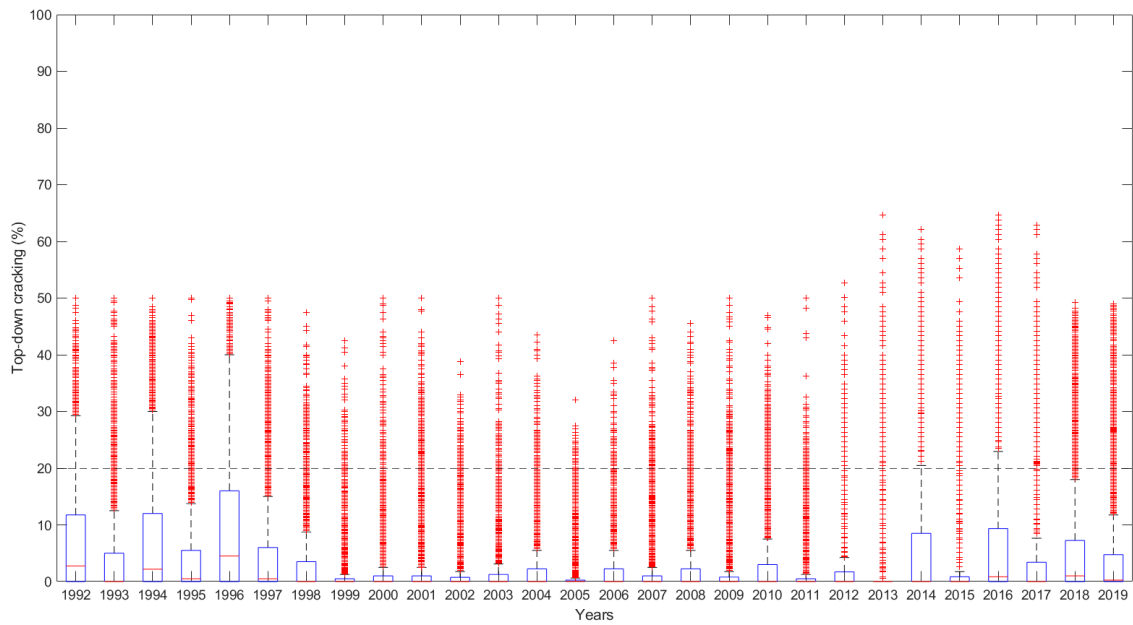


Figure 3-10 Top-down cracking at every 0.1-mile segment for flexible sections

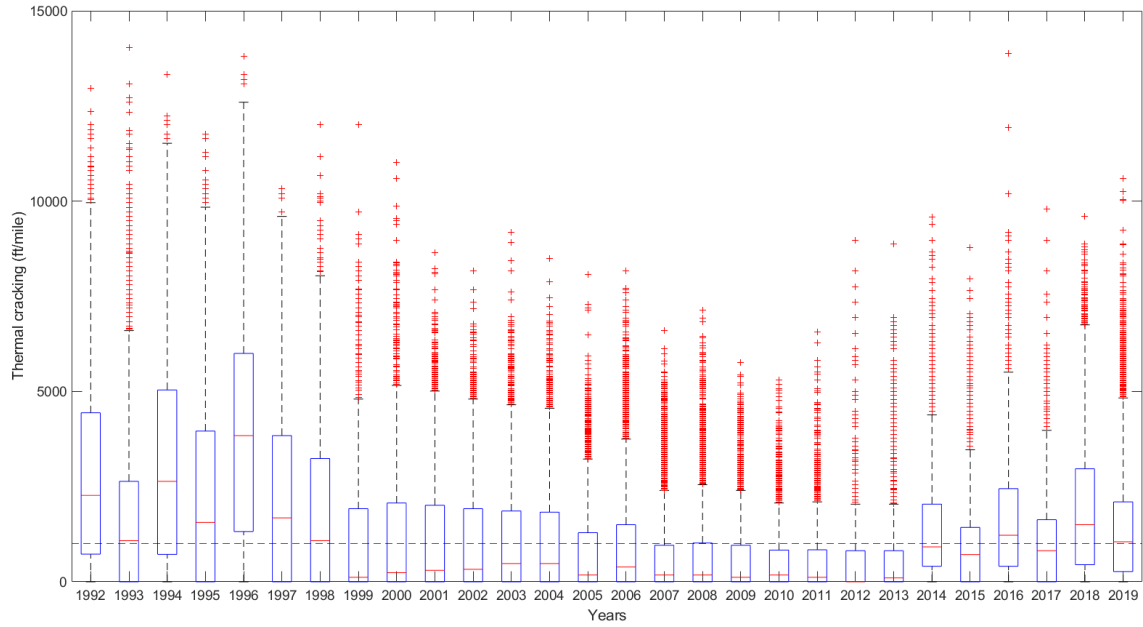


Figure 3-11 Thermal cracking at every 0.1-mile segment for flexible sections

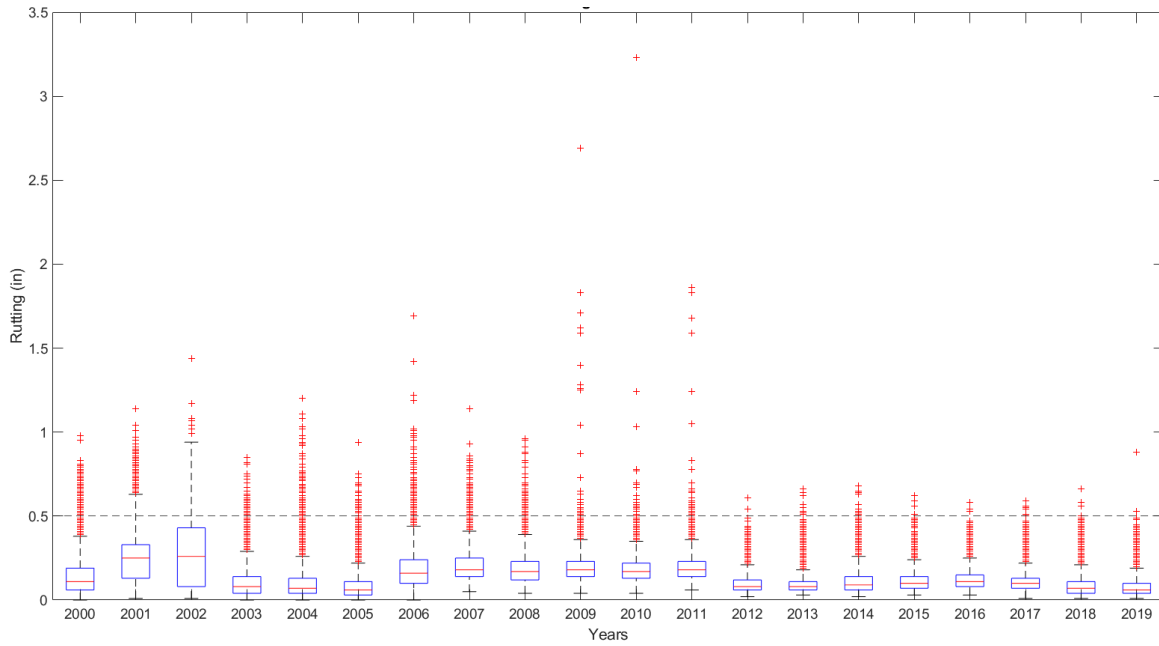


Figure 3-12 Rutting at every 0.1-mile segment for flexible sections

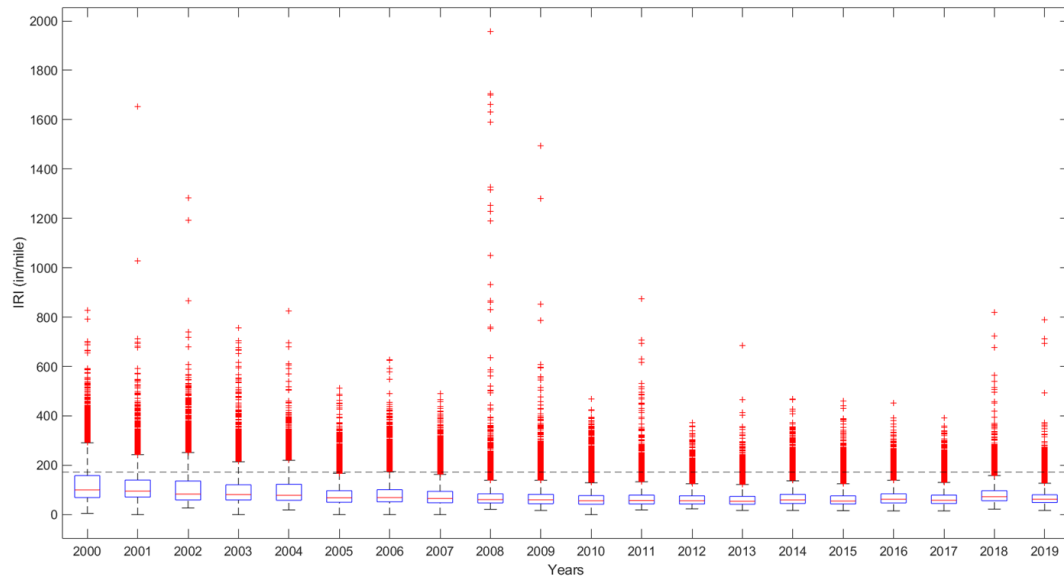


Figure 3-13 IRI at every 0.1-mile segment for flexible sections

Figures 3-14 to 3-16 show the raw performance data for all rigid sections. As previously noted, 2112 ft/mile and 0.4 inches cutoff values were adopted for thermal cracking and joint faulting, respectively. These values were selected based on the raw (0.1-mile segment) data, limitations of the Pavement-ME models, and consensus with MDOT. Moreover, sections with Superpave mixes are only used to calibrate the thermal cracking model to have consistent Level 1 input in the Pavement-ME.

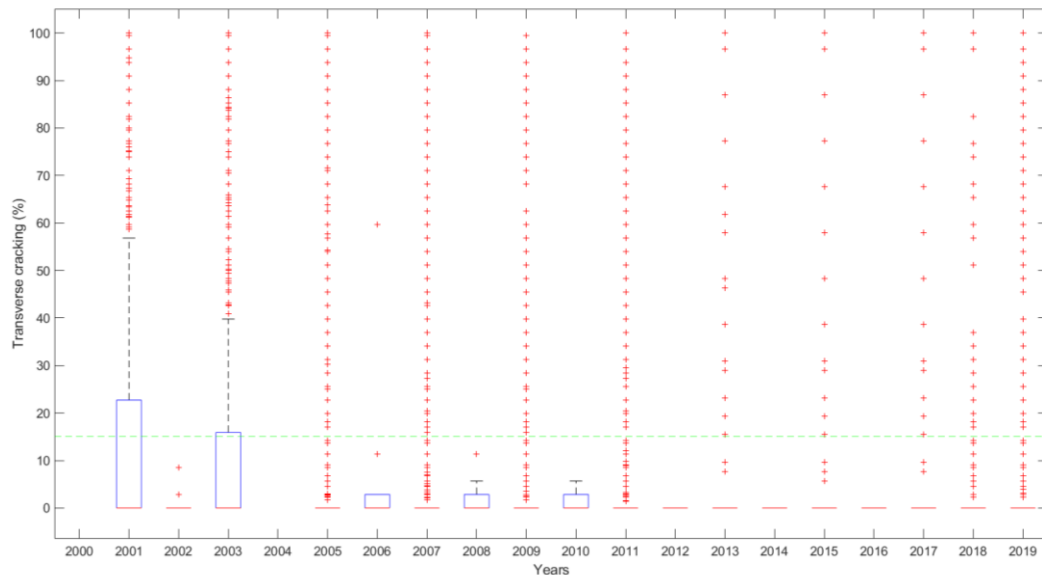


Figure 3-14 Transverse cracking at every 0.1-mile segment for rigid sections

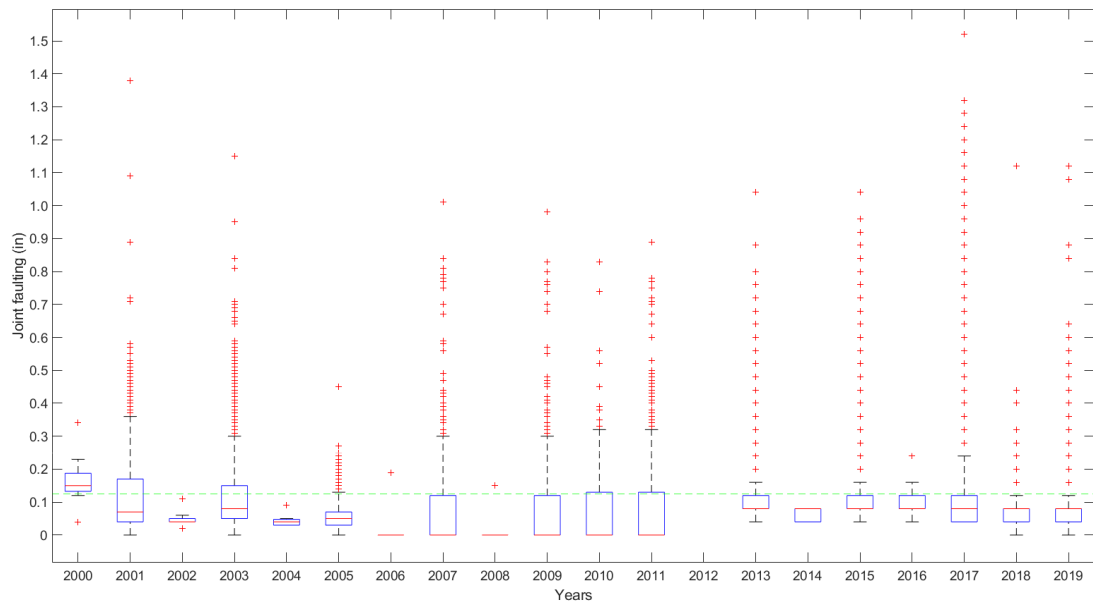


Figure 3-15 Joint faulting at every 0.1-mile segment for rigid sections

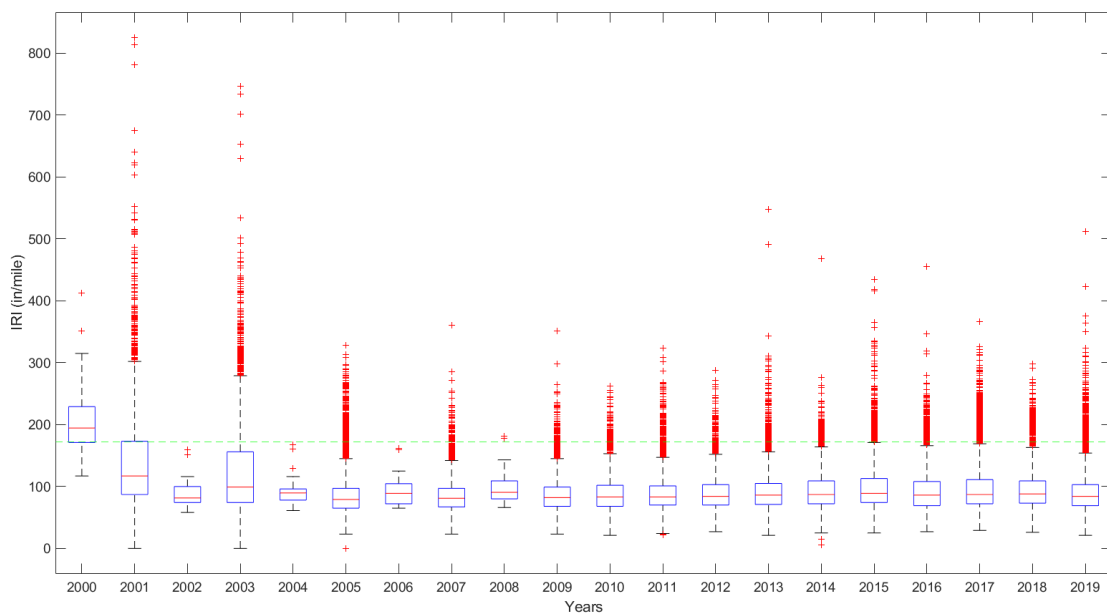


Figure 3-16 IRI at every 0.1-mile segment for rigid sections

3.4 SELECTED SECTION PERFORMANCE DATA SUMMARY

The measured performance data was extracted for each project, and the necessary conversions were made to ensure compatibility with the Pavement-ME predicted performance, as discussed in Section 3.2. The level of distress was assessed in all pavement sections identified for local

calibration. The calibration process entails comparing each chosen project's predicted and measured performance. To have a robust local calibration, the levels of distress must fall within a reasonable range (i.e., above and below threshold limits for each type of distress). Therefore, the distress levels for all projects were compiled and analyzed to determine their respective ranges. This section summarizes the observed performance for the selected flexible and rigid pavement sections. Efforts were undertaken to gather sufficient information to achieve a precise and dependable local calibration of the performance models. Due to changes in construction practices and/or data availability, most sections are less than 20 years old, so it is expected that most sections do not have poor performance or exceed performance thresholds. Furthermore, these represent the average values of the Pavement-ME prediction using 50% reliability. When designing, a higher reliability factor is applied to account for project variability (including climate, traffic, material, and construction), increasing the resulting distress values. Therefore, while designs will correlate with the calibration sections, it should not be anticipated that pavement designs will exactly match the sections used in calibration because of the increased reliability factor.

3.4.1 Flexible Performance Data

The magnitude and age distribution for the HMA reconstruct sections (also includes crush and shape and HMA over rubblized PCC) are shown in Figures 3-17 to 3-21. The following observations were made:

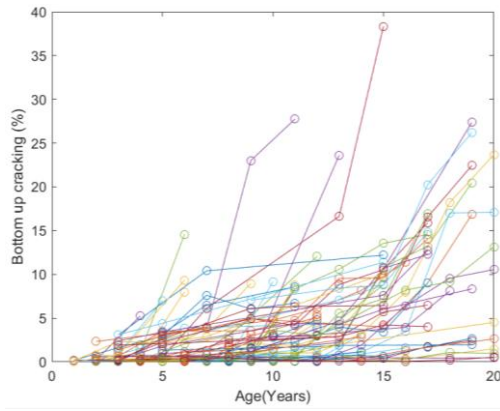
- Bottom-up cracking: Bottom-up cracking magnitudes are usually low for most sections, with only a seven crossing the threshold of 20% with a maximum of almost 40%. The maximum age ranges from 4 to 20 years. Most sections fall in the good category, as shown in Figure 3-8.
- Longitudinal/top-down cracking: Top-down cracking is observed more frequently than bottom-up cracking. More sections have observed top-down cracking compared to bottom-up cracking. The age at maximum distress ranged from 5 to 20 years.
- Thermal cracking: Higher thermal cracking values are observed, ranging up to 4000 ft/mile. The design threshold used by MDOT is 1000 ft/mile. The age at which the maximum thermal cracking is observed ranges from 5 to 19 years. Sections with performance grade (PG) binders have been used for thermal cracking calibration.

- Rutting: Selected sections do not exhibit significant rutting. All sections were below the threshold of 0.5 inches. The age distribution ranged from 3 to 19 years. Two-thirds of the sections are in the fair performance category, as shown in Figure 3-8.
- IRI: The IRI time series is usually flat, with no sections exceeding the 172 in/mile threshold. The maximum observed IRI is 168.5 in/mile. The age at maximum IRI ranged from 5 to 20 years. It is worth noting that a cutoff value of the initial IRI less than or equal to 77 in/mile is selected to calibrate the IRI model. 74% of sections are in good, followed by 25% of sections in fair category. Only 1% of sections showed poor performance.

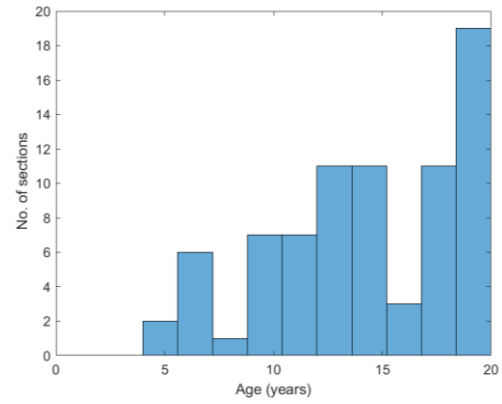
3.4.2 Rigid Performance Data

The magnitude and age distribution for the JPCP rehabilitation projects are shown in Figures 3-22 to 3-24. The following observations can be made from the figures:

- Transverse cracking: A maximum transverse cracking value of 85% is observed, with five sections crossing the distress threshold of 15% slabs cracked. The age distribution ranges from 4 to 20 years. About 72% of these sections fall under the fair performance category, as shown in Figure 3-7.
- Transverse joint faulting: Ten sections exceed the joint faulting threshold of 0.125 inches, with a maximum value of 0.17 inches. The age distribution ranges from 8 to 20 years. These observed values for joint faulting have been cut off at 0.4 inches, where a 0.1-mile segment is above 0.4 inches.
- IRI: A maximum IRI of 167 in/mile was observed. The age at maximum IRI ranges from 5 to 20 years. It is worth noting that a cutoff value for the initial IRI less than or equal to 82 in/mile is used to calibrate the IRI model. All sections fall under good and fair categories, with none exhibiting poor performance, as shown in Figure 3-7.

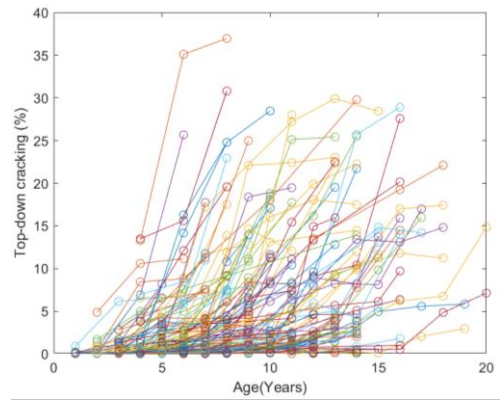


(a) Time series

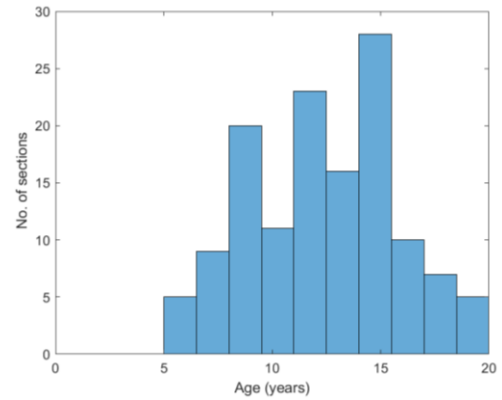


(b) Age distribution

Figure 3-17 Selected flexible sections — Bottom-up cracking

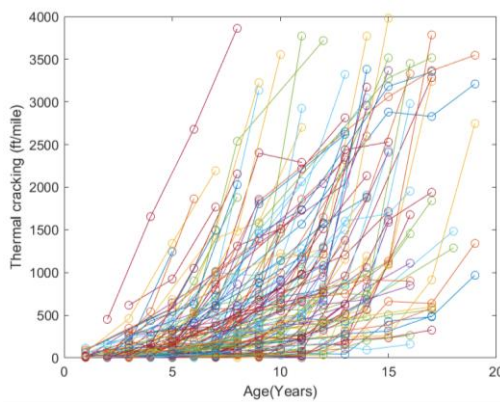


(a) Time series

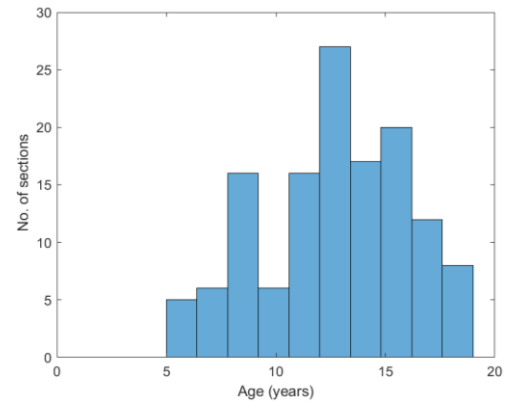


(b) Age distribution

Figure 3-18 Selected flexible sections — Top-down cracking

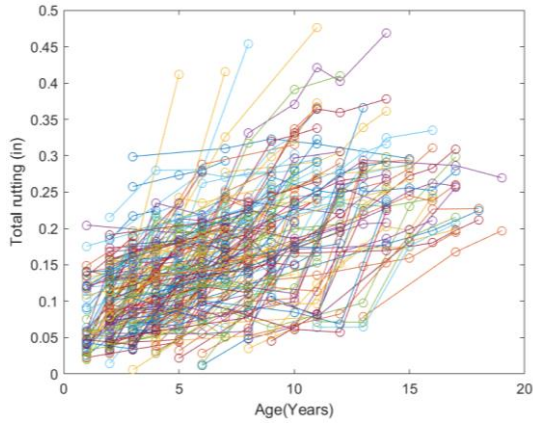


(a) Time series

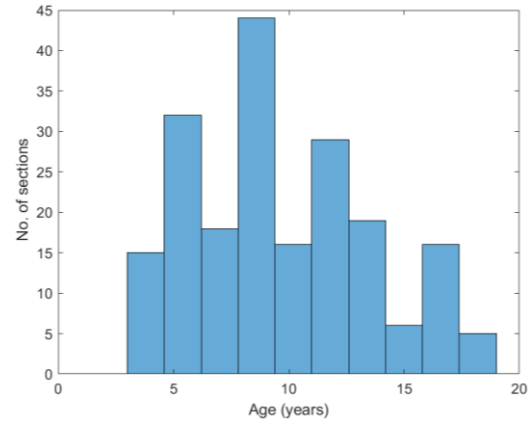


(b) Age distribution

Figure 3-19 Selected flexible sections— Transverse (thermal) cracking

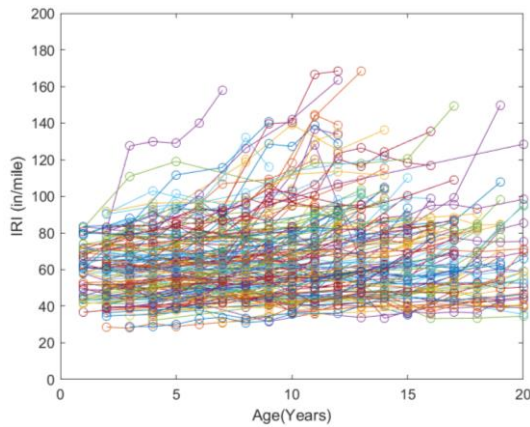


(a) Time series

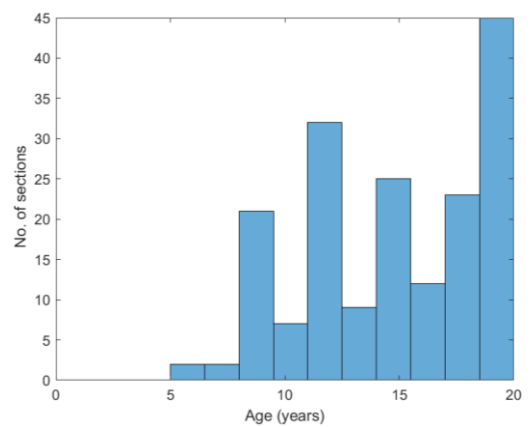


(b) Age distribution

Figure 3-20 Selected flexible sections — Total rutting

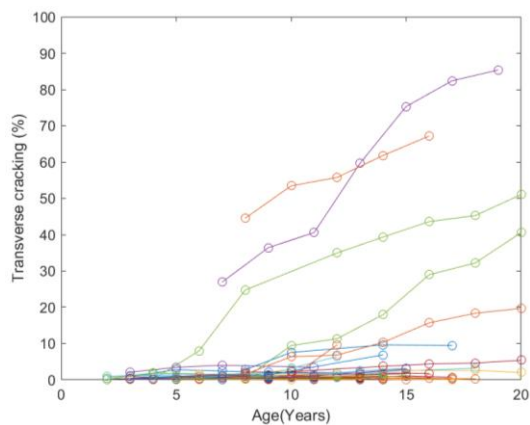


(a) Time series

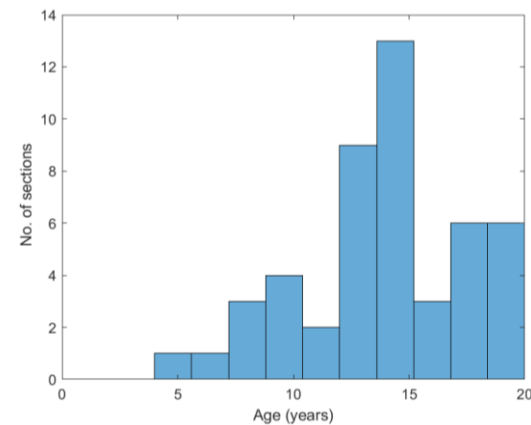


(b) Age distribution

Figure 3-21 Selected flexible sections — IRI



(a) Time series



(b) Age distribution

Figure 3-22 Selected rigid sections — Transverse cracking

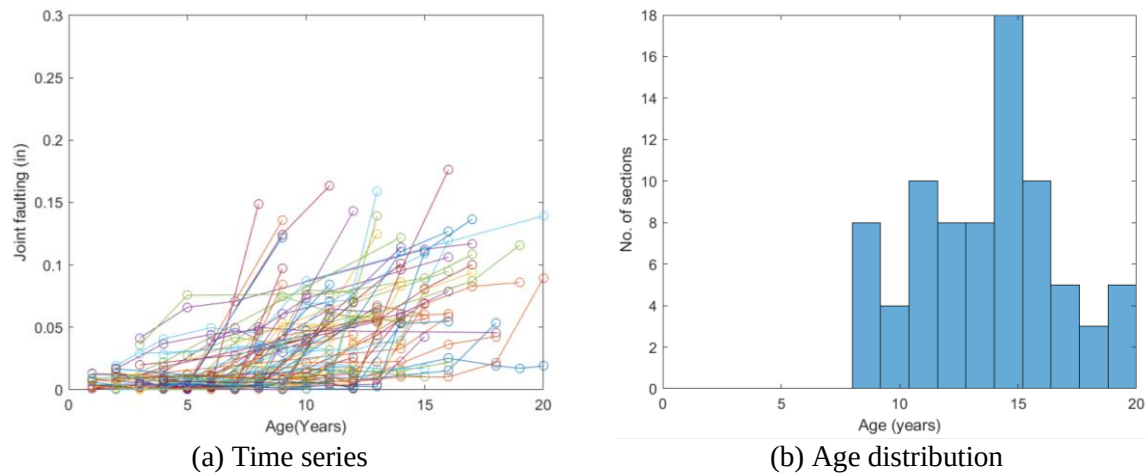


Figure 3-23 Selected rigid sections — Joint faulting

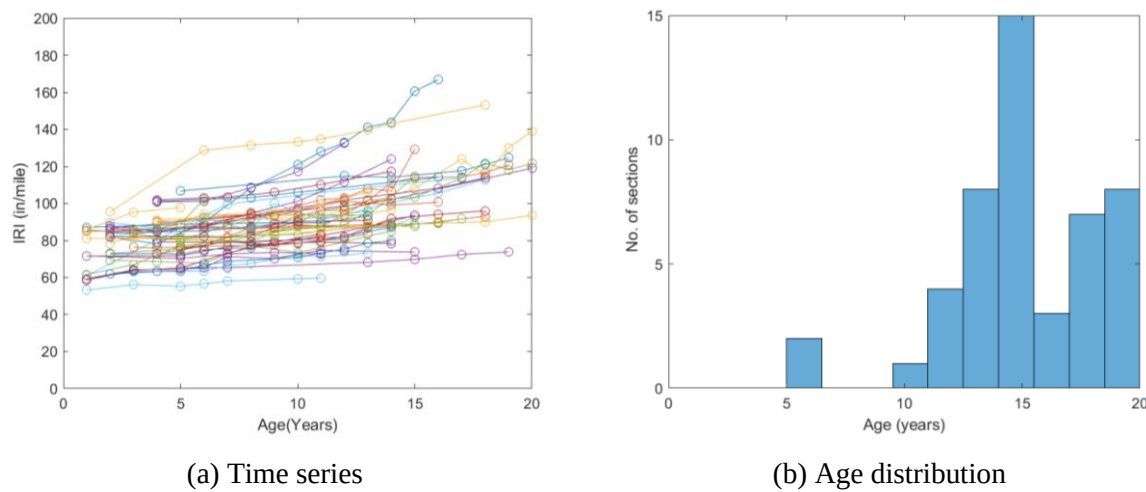


Figure 3-24 Selected rigid sections — IRI

3.5 INPUT DATA EXTENT

Accurate pavement cross-sectional, traffic, climate, and material input data are essential for adequately characterizing as-constructed pavements since the information directly affects performance prediction accuracy in the Pavement-ME software. Due to the large number of inputs required to characterize a pavement in the Pavement-ME, input data collection can be time-consuming. Moreover, many critical input parameters have three input levels within the Pavement-ME hierarchical structure. The process of collecting as-constructed input data, including the source of the data, how to address missing data, and the selection of input values, is discussed in this section. The best available input level was used for the selected pavement sections.

3.5.1 Pavement Cross-Section

The pavement cross-sectional information is necessary to characterize the layer thicknesses of the various layers. The cross-sectional information is obtained from the construction records. Typically, in the case of HMA pavements, the drawings provided the asphalt application rate of the HMA layers (dividing the application rate by 110), which was used to determine the HMA lift thicknesses in inches. For the sections used in the previous calibration effort (10), the Pavement-ME inputs data sheet was used to extract design inputs. MDOT provided the drawings (construction plans) for the newly selected sections. The thickness, mix type, traffic, and unbound layer information were included in these drawings. A summary of the design thicknesses for flexible and rigid selected pavement projects is shown in Tables 3-6 and 3-7.

Table 3-6 Average flexible pavement thicknesses

Pavement types	HMA top course thickness (in.)	HMA leveling course thickness (in.)	HMA base course thickness (in.)	Base thickness (in.)	Subbase thickness (in.)
Crush and Shape	1.6	1.9	2.0	7.5	20.5
Freeway	1.6	2.1	4.5	7.1	16.8
Non-freeway	1.5	2.1	3.2	6.6	16.4
Rubblized	1.6	2.0	3.0	3.8	11.1
Statewide Average	1.6	2.0	3.1	5.7	15.0

Table 3-7 Average rigid pavement thicknesses

Pavement type	Average PCC thickness (in.)	Average base thickness (in.)	Average subbase thickness (in.)
JPCP	11.4	6.9	12.1

3.5.2 Traffic Inputs

The traffic data is a critical input to the Pavement-ME. Level 2 traffic data was used for all sections. MDOT provided a spreadsheet with traffic distribution tables, which was used to extract Pavement-ME inputs for traffic. These tables include:

- Vehicle class distribution
- Hourly distribution (only for rigid sections)
- Monthly adjustment factor
- Number of axles per truck
- Single axle load spectra

- Tandem axle load spectra
- Tridem axle load spectra
- Quad axle load spectra

The inputs (with input categories) required to obtain these tables are summarized in Table 3-8.

Table 3-8 Traffic input categories

Inputs	Categories
Percentage of vehicle class 9	• Less than 45 • 45 to 70 • Above 70
Region type	• Rural • Urban
COHS type	• National • Regional • Statewide
Number of lanes	• 2 • 3 • 4+

The number of lanes was identified from the plans. Wherever the number of lanes was unavailable, they were visually estimated utilizing Google Maps coordinates. The COHS (Corridors of Highest Significance) type was estimated using each project's PR number and beginning and ending milepost. The percentage of class 9 vehicles was estimated for each section using the MDOT Transportation Data Management System (TDMS) website from the following URL: <https://mdot.public.ms2soft.com/tcds/tsearch.asp?loc=mdot>. For sections where the traffic data was unavailable at the exact location, nearby locations in the same section were used. The range and average two-way AADTT values for all flexible and rigid sections are summarized in Table 3-9.

Table 3-9 Ranges of AADTT for all reconstruct projects

Road Type	Min AADTT	Max AADTT	Average AADTT
Crush and Shape	60	1986	669
Rubblized	173	3707	1502
HMA Reconstruct (Freeway)	313	6745	2076
HMA Reconstruct (Non-freeway)	63	1600	431
JPCP Reconstruct	150	18297	7141
Statewide Average	134	6502	2381

3.5.3 As-constructed Material Inputs

The as-constructed material inputs were obtained from the construction records, JMFs, and other test records. Ideally, these inputs are to be recorded at the time of construction. These inputs range between project-specific and statewide average values. This section details the material properties of each pavement structural layer.

3.5.3.1. HMA layer inputs

All inputs were collected at the highest hierarchy level; however, the needed data were unavailable for all pavement sections. In that case, the data was collected using other correlations/sources. Data collection for each HMA layer input is as follows:

- Dynamic modulus (E^*): E^* was obtained from the DYNAMOD software developed in a previous study (71). E^* for the Superpave mixes was directly obtained from the database. For older mixes (marshal mixes), the volumetric, binder, and gradation information was used to predict the E^* using DYNAMOD's Artificial Neural Networks (ANNs). E^* was obtained at Level 1.
- Binder (G^*): G^* was also obtained from the DYNAMOD database using the region and binder information. G^* was obtained at Level 1.
- Creep compliance ($D(t)$): $D(t)$ was obtained from the DYNAMOD database. $D(t)$ was obtained at Level 1 for Performance grade (PG) sections and Level 3 for other sections.
- Indirect tensile strength (IDT): IDT was obtained from the DYNAMOD database at Level 2 for Performance grade (PG) sections and Level 3 for other sections.
- AC layer thickness: These were obtained from construction records. Usually, the application rate in lbs/yards² is available, which can be utilized to obtain the layer thickness, as previously mentioned.
- Air voids and binder content: As constructed air voids and binder content were obtained from construction records. Table 3-10 summarizes the average as-constructed air voids for different pavement types. Historical test records were utilized for unavailable data to obtain an average value based on mix type, as shown in Table 3-12.
- Aggregate gradation: Gradation was obtained from JMFs. Tables 3-11 summarize the average gradation for the top, leveling, and base layers, respectively, for different

pavement types. Historical test records were utilized for unavailable data to obtain an average value based on mix type, as shown in Table 3-12.

It is important to note that Level 1 G^* and Level 2 IDT data were used to calibrate the thermal cracking model.

Table 3-10 As-constructed percent air voids for HMA layers

HMA layer	Road Type	Average as-constructed air voids
Top course	Crush and Shape	6.1
	Rubblized	6.8
	HMA Reconstruct Freeway	6.6
	HMA Reconstruct Non-freeway	6.8
Leveling course	Crush and Shape	6.2
	Rubblized	6.4
	HMA Reconstruct Freeway	6.7
	HMA Reconstruct Non-freeway	6.7
Base course	Crush and Shape	5.8
	Rubblized	5.8
	HMA Reconstruct Freeway	6.4
	HMA Reconstruct Non-freeway	6.8

Table 3-11 HMA layer average aggregate gradation

HMA Layer	Road type	Effective AC binder content	Percent passing sieve size			
			3/4	3/8	#4	#200
Top course	Crush and Shape	11.5	100.0	89.7	68.4	5.2
	Rubblized	11.9	99.4	89.8	67.3	5.9
	HMA Reconstruct (Freeway)	11.2	100.0	92.4	67.4	5.2
	HMA Reconstruct (Non-freeway)	11.1	100.0	94.6	71.4	5.3
Leveling course	Crush and Shape	10.6	100.0	81.8	61.1	5.0
	Rubblized	11.2	100.0	87.0	67.8	5.2
	HMA Reconstruct (Freeway)	10.1	99.8	81.3	63.3	4.8
	HMA Reconstruct (Non-freeway)	10.2	100.0	82.6	73.4	4.8
Base course	Crush and Shape	10.8	99.6	77.9	60.3	4.6
	Rubblized	10.6	99.3	78.9	59.9	4.8
	HMA Reconstruct (Freeway)	9.4	95.8	72.9	51.6	4.9
	HMA Reconstruct (Non-freeway)	9.6	98.9	76.6	57.5	4.9

Table 3-12 MDOT recommended values volumetrics and gradation

Mix type	Air voids (%)	Effective binder content (%)	% Passing 3/4" Sieve	% Passing 3/8" Sieve	% Passing #4 Sieve	% Passing #200 Sieve
3E1	5.8	10.8	99.85	80.44	62.94	4.40
4E1	6.1	11.5	100.00	87.24	70.43	5.11
5E1	6	12.6	100.00	97.14	78.23	5.63
2E3	4.8	9.7	92.65	68.70	53.95	4.40
3E3	5.8	10.8	99.63	77.88	60.33	4.56
4E3	6.1	11.5	100.00	86.91	68.66	4.92
5E3	6	12.6	100.00	97.86	79.81	5.49
2E10	4.8	9.7	94.55	73.50	59.70	4.50
3E10	5.8	10.8	99.78	80.27	62.78	4.84
4E10	6.1	11.5	100.00	87.65	70.06	5.26
5E10	6	12.6	100.00	98.30	81.27	5.67
2E30	4.8	9.7	99.00	71.80	60.60	4.20
3E30	5.8	10.8	99.95	79.20	59.82	4.40
4E30	6.1	11.5	100.00	88.63	66.90	4.33
5E30	6	12.6	100.00	99.00	81.24	5.68

3.5.3.2. PCC material inputs

The Pavement-ME transverse cracking prediction model is very sensitive to concrete strength (compressive or flexural). The PCC material-related inputs were obtained from material testing results. If these results were unavailable, typical MDOT values were used.

PCC strength:

MDOT collected the concrete core compressive strength (f'_c) test data. These tests represent the concrete compressive strength close to the construction time for the selected pavement sections. These test values were used directly for each corresponding project. If compressive strength is unavailable, an average value of 5239 psi was used. This is an average value obtained from the sections with available values. The transverse cracking model in the Pavement-ME directly uses the modulus of rupture (MOR) to estimate the damage. The MOR values were calculated based on the ACI correlation between MOR and f'_c (used in the Pavement-ME), as shown by Equation (3-7). Figure 3-25 shows the f'_c and estimated MOR distributions. It should be noted that these cores' specific testing age was unavailable; however, all cores were tested after or at least 28 days. The Pavement-ME internally calculates the relationship between f'_c and MOR.

$$MOR = 9.5 \times \sqrt{f'_c} \quad (3-7)$$

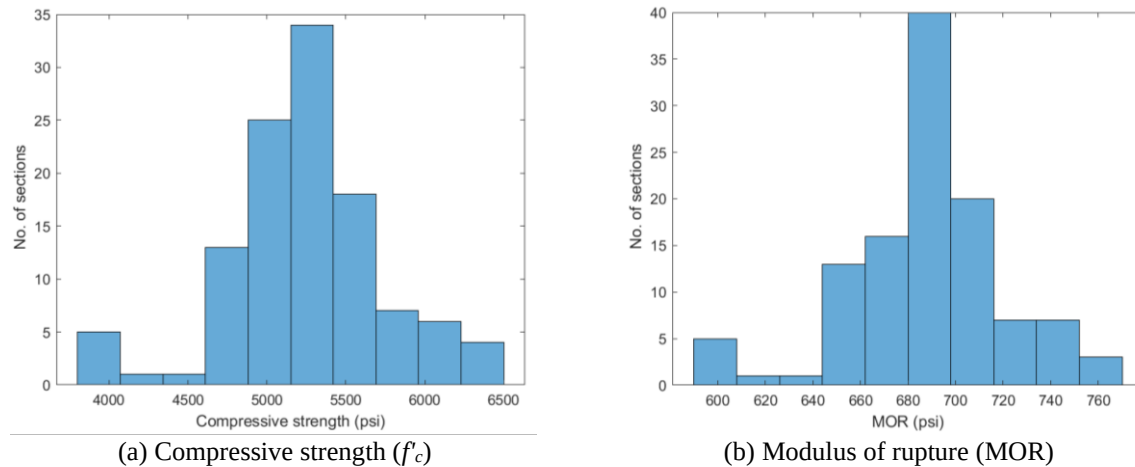


Figure 3-25 Distribution of concrete strength properties

Coefficient of thermal expansion:

The CTE input values were obtained from the MDOT recommended values (72). A value of $4.4 \text{ in/in/}^\circ\text{F} \times 10^{-6}$ was used for Bay, Grand, North, Southwest, and Superior regions, whereas $5.0 \text{ in/in/}^\circ\text{F} \times 10^{-6}$ was used for Metro and University regions.

3.5.3.3. Aggregate base/subbase and subgrade input values

The aggregate base/subbase and subgrade input values were obtained from the following sources:

- Backcalculation of unbound granular layer moduli (73)
- Pavement subgrade MR design values for Michigan's seasonal changes (74)

The resilient modulus (MR) values for the base and subbase material were selected based on the results from previous MDOT studies (73, 74). The typical backcalculated values for base and subbase MR is 33,000 psi and 20,000 psi, respectively. It is worth noting that crushed and shaped and rubblized sections have been modeled as new flexible pavements. The existing layer has been modeled as a dense aggregate base with an MR of 125,000 psi for crush and shape and 70,000 for rubblized sections. These values were assumed to be the same for all projects since in-situ MR values were unavailable. For base/subbase layers, the software default to "Modify input values by temperature/moisture" was selected. The subgrade material type and resilient modulus were selected based on the Subgrade MR study (73, 74). The study outlined the location of specific soil types and their MR values across the entire State. Annual representative values for

subgrade MR were used in Pavement-ME. The recommended design MR value corresponding to the soil type is shown in Table 3-13.

Table 3-13 Average roadbed soil MR values

Roadbed Type		Average MR		
USCS	AASHTO	Laboratory-determined (psi)	Back-calculated (psi)	Recommended design MR value (psi)
SM	A-2-4, A-4	17,028	24,764	5,200
SP1	A-1-a, A-3	28,942	27,739	7,000
SP2	A-1-b, A-3	25,685	25,113	6,500
SP-SM	A-1-b, A-2-4, A-3	21,147	20,400	7,000
SC-SM	A-2-4, A-4	23,258	20,314	5,000
SC	A-2-6, A-6, A-7-6	18,756	21,647	4,400
CL	A-4, A-6, A-7-6	37,225	15,176	4,400
ML	A-4	24,578	15,976	4,400
SC/CL/ML	A-2-6, A-4, A-6, A-7-6	26,853	17,600	4,400

3.5.4 Climatic Inputs

The Enhanced Integrated Climatic Model (EICM) in Pavement-ME requires hourly climatic data. This data includes air temperature, precipitation, relative humidity, percent sunshine, and wind speed. A statistical comparison between Modern-Era Retrospective Analysis for Research and Applications (MERRA) and North American Regional Reanalysis (NARR) data was performed to identify the most suitable climatic data for calibration. Both MERRA and NARR data files used include climatic information for different periods. For that purpose, a common temporal overlap of 13 years was identified for which continuous hourly data is available for all climatic files from September 2000 to September 2013. The MERRA stations falling in the lake region were removed from the database. Moreover, the four closest MERRA stations were identified for each NARR station, and the weighted average (proportional to the distance) for all four stations based on their distances was used for comparison. A total of 29 NARR stations and the four closest corresponding MERRA stations to each have been compared. Table 3-14 shows the SEE, bias, and correlation coefficient (R) between MERRA and NARR for hourly, daily, and monthly data (75). MDOT has been using default Pavement-ME climate data and ground-based climate automated surface observation systems (ASOS) data. This data was reviewed for

errors/anomalies and was improved in MDOT's previous study (68). The following observations were made based on the comparison and previous study (68, 75):

- MERRA and NARR climatic data are comparable for air temperature followed by humidity and wind speed. Percent sunshine showed a low correlation, and precipitation data is significantly different (i.e., a very low correlation) among all climatic inputs.
- The predicted pavement performance using MERRA-2 and NARR climatic data showed good agreement except for thermal cracking in flexible pavement and transverse cracking in rigid pavements. These differences are expected mainly because of sunshine data.
- MERRA has anomalies in humidity data. Several humidity values were erroneously higher than 100.
- MERRA appeared to be incorrectly estimating precipitation. Specifically, the number of wet days was extremely high, such that the data review showed wet event days in the data on actual dry days. The ground-based stations are more closely aligned with actual wet event days. Furthermore, it was unclear why the percent sunshine was significantly different.

Table 3-14 Descriptive statistics for MERRA and NARR data comparison

Climatic input	Descriptive statistics	Hourly			Daily			Monthly		
		SEE	Bias	R	SEE	Bias	R	SEE	Bias	R
Humidity	Mean	12.784	4.437	0.764	9.582	4.437	0.705	7.387	4.437	0.538
	Std. Dev.	0.726	2.230	0.035	1.014	2.230	0.055	1.283	2.230	0.145
	COV	5.68%	50.27%	4.60%	10.58%	50.27%	7.86%	17.37%	50.27%	26.96%
Precipitation	Mean	0.049	0.002	0.062	0.009	0.002	0.610	0.002	0.002	0.678
	Std. Dev.	0.005	0.000	0.022	0.001	0.000	0.045	0.000	0.000	0.059
	COV	10.85%	15.22%	34.59%	7.90%	15.22%	7.33%	11.21%	15.22%	8.73%
Sunshine	Mean	44.614	-1.457	0.411	29.317	-1.457	0.570	11.847	-1.457	0.821
	Std. Dev.	3.908	6.809	0.071	2.777	6.809	0.079	1.788	6.809	0.033
	COV	8.76%	-467.39%	17.27%	9.47%	-467.39%	13.84%	15.09%	-467.39%	4.04%
Temperature	Mean	3.924	-0.771	0.982	2.710	-0.771	0.992	1.837	-0.771	0.997
	Std. Dev.	0.548	0.766	0.006	0.436	0.766	0.003	0.428	0.766	0.002
	COV	13.98%	-99.43%	0.58%	16.08%	-99.43%	0.31%	23.32%	-99.43%	0.20%
Wind speed	Mean	3.318	-0.165	0.752	2.031	-0.165	0.863	1.470	-0.165	0.848
	Std. Dev.	0.946	1.700	0.100	1.097	1.700	0.105	1.145	1.700	0.145
	COV	28.52%	-1029.25%	13.25%	54.00%	-1029.25%	12.16%	77.92%	-1029.25%	17.10%

Note: $SSE = \sqrt{\frac{\sum(MERRA-NARR)^2}{n-2}}$; $Bias = \frac{\sum(MERRA-NARR)}{n}$

In the previous study, additional weather stations were added to improve the climate coverage using ASOS and the Michigan Road Weather Information System (RWIS) as potential data sources (68). Moreover, additional years of climatic data were added from February 2006 to December 2014 to enhance the data. Since the predicted performance did not show significant differences and the NARR data was improved for Michigan climate, the improved MDOT

NARR climatic files were used for climatic inputs for both flexible and rigid pavements. The files were downloaded as *.hcd files, which can be read directly in Pavement-ME. The closest weather station to each selected project was used.

Table 3-15 Michigan climate station information

HCD filename	City/Location	Climate identifier	Latitude	Longitude
4847	Adrian	Adrian Lenawee County Arpt	41.868	-84.079
94849	Alpena	Alpena Co Rgnl Airport	45.072	-83.581
94889	Ann Arbor	Ann Arbor Municipal Arpt	42.224	-83.74
14815	Battle Creek	W K Kellogg Airport	42.308	-85.251
94871	Benton Harbor	Sw Michigan Regional Arpt	42.129	-86.422
14822	Detroit	Detroit City Airport	42.409	-83.01
94847	Detroit	Detroit Metro Wayne Co Apt	42.215	-83.349
14853	Detroit	Willow Run Airport	42.237	-83.526
14826	Flint	Bishop International Arpt	42.967	-83.749
4854	Gaylord	Otsego County Airport	45.013	-84.701
94860	Grand Rapids	Gerald R Ford Intl Airport	42.882	-85.523
14858	Hancock	Houghton County Memo Arpt	47.169	-88.506
4839	Holland	Tulip City Airport	42.746	-86.097
94814	Houghton Lake	Roscommon County Airport	44.368	-84.691
94893	Iron Mountain/Kingsford	Ford Airport	45.818	-88.114
14833	Jackson	Jakson Co-Rynolds Fld Arpt	42.26	-84.459
94815	Kalamazoo	Klmazo/Btl Creek Intl Arpt	42.235	-85.552
14836	Lansing	Capital City Airport	42.78	-84.579
14840	Muskegon	Muskegon County Airport	43.171	-86.237
14841	Pellston	Pton Rgl Ap Of Emmet Co Ap	45.571	-84.796
94817	Pontiac	Oakland Co. Intl Airport	42.665	-83.418
14845	Saginaw	Mbs International Airport	43.533	-84.08
14847	Sault Ste Marie	Su Ste Mre Muni/Sasn Fl Ap	46.467	-84.367
14850	Traverse City	Cherry Capital Airport	44.741	-85.583
AMN	Alma	Gratiot Community Airport	43.322	-84.688
BAX	Bad Axe	Huron County Memorial Airport	43.78	-82.985
CFS	Caro	Tuscola Area Airport	43.459	-83.445
ERY	Newberry	Luce County Airport	46.311	-85.4572
ESC	Escanaba	Delta County Airport	45.723	-87.094
FKS	Frankfort	Frankfort Dow Memorial Field Airport	44.625	-86.201
IRS	Sturgis	Kirsch Municipal Airport	41.813	-85.439
ISQ	Manistique	Schoolcraft County Airport	45.975	-86.172
IWD	Ironwood	Gogebic Iron County Airport	46.527	-90.131
LDM	Ludington	Mason County Airport	43.962	-86.408
MOP	Mount Pleasant	Mount Pleasant Municipal Airport	43.622	-84.737
OSC	Oscoda	Oscoda Wurtsmith Airport	44.452	-83.394
PHN	Port Huron	Saint Clair County Intl Airport	42.911	-82.529
RQB	Big Rapids	Roben Hood Airport	43.723	-85.504
SAW	Gwinn	Sawyer International Airport	46.354	-87.39

These files were directly used for rigid sections (since they are default files in the Pavement-ME), and custom stations were formed using these files for flexible sections. Table 3-15 summarizes the climatic files used for calibration.

3.5.5 Estimation of Initial IRI

Initial IRI is an essential input for IRI prediction and pavement design. Initial IRI is the IRI value right after the construction. It indicates construction and ride quality right after construction. Initial IRI is also an essential part of QC/QA testing. Moreover, higher initial IRI values may lead to a reduction in pavement service life. The IRI model in the Pavement-ME is linear in form, but the measured IRI data may not always be linear. The change in measured IRI with time can be linearly increasing or non-linearly increasing, which may follow an irregular or flat trend. Also, the initial IRI (if available) can be greater or smaller than the first measured IRI data points because of the measurement date. Figure 3-26 shows some examples of measured IRI trends for flexible and rigid sections.

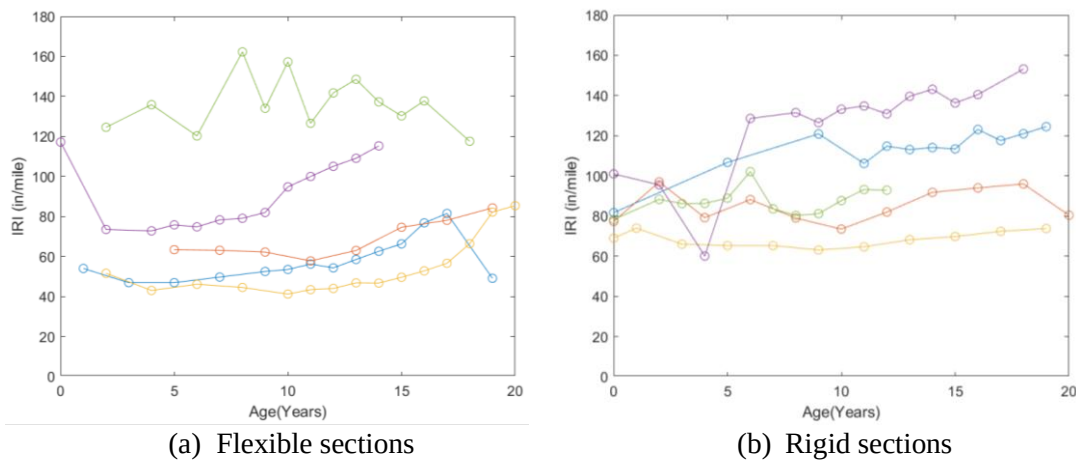


Figure 3-26 Examples of measured IRI trends

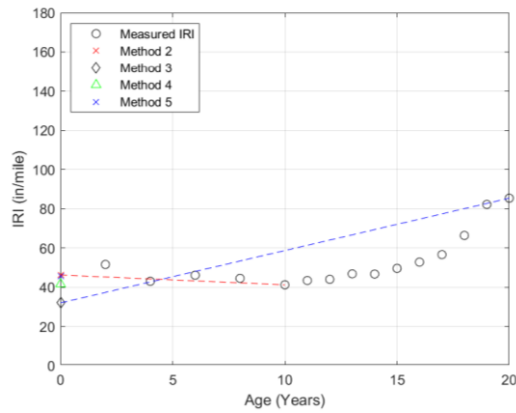
A single backcasting approach may not be applicable for all sections due to the difference in measured IRI trends for each section. Considering the data limitations and challenges, a systematic approach is used to estimate the initial IRI. Five different methods used include:

1. Selecting the IRI at zeroth year (if available).
2. Linear backcasting IRI based on the measured data for the first ten years.
3. Linear backcasting IRI based on the measured data for all available years.
4. Reducing the first measured IRI (after construction) by 5 inches per mile/year up to the zeroth year.
5. Reducing the first measured IRI (after construction) by 5 inches per mile/year if greater than 100; 4 inches per mile/year if between 70 and 100; 3 inches per mile/year if less than 70 up to a zeroth year.

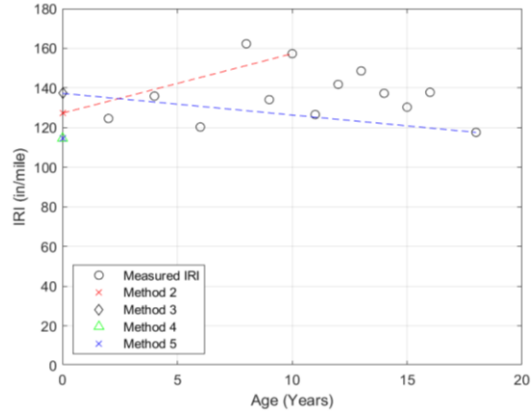
It is important to note that the MDOT specification limit of 70 in/mile and 75 in/mile for flexible and rigid pavements are considered. After the initial IRI was obtained using the five methods mentioned above, the final initial IRI was selected based on the following criteria:

1. Use the initial IRI (if available) if it is less than or equal to the specification limit.
2. If the initial IRI (if available) is greater than the specification limit, use the backcasted IRI from other methods, whichever is closest to and lower than the specification limit.
3. If all five methods provide an initial IRI greater than the specification limit, choose the approach with an initial IRI greater than and closest to the specification limit.
4. Subsequently, review data progression to see if the estimated initial IRI fits all available measured data points.

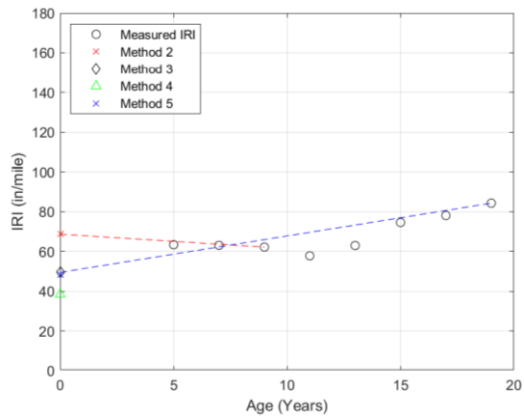
Figure 3-27 shows example sections with backcasted initial IRI using different methods. Section 1 has a non-linearly increasing trend, section 2 has an irregular trend, and sections 3 and 4 have linear trends with varying slopes. Different backcasting methods provide significantly different initial IRI values. For example, section 2 has a maximum difference of more than 20 inches/mile among the initial IRI values calculated using various methods. Similar differences can be seen in other sections. Moreover, method 3 for section 2 provides an unrealistic initial IRI value, higher than the first measured data point, due to the nature of the irregular trend. These plots show a need for different backcasting methods for various IRI trends. Figure 3-28 shows a flowchart for selecting the initial IRI using the mentioned approach. For certain flexible sections, the final selected initial IRI was very low (less than 30 in/mile). In that case, an initial IRI value of 30 in/mile is assigned. Moreover, the final selected initial IRI was very high for several flexible and rigid sections.



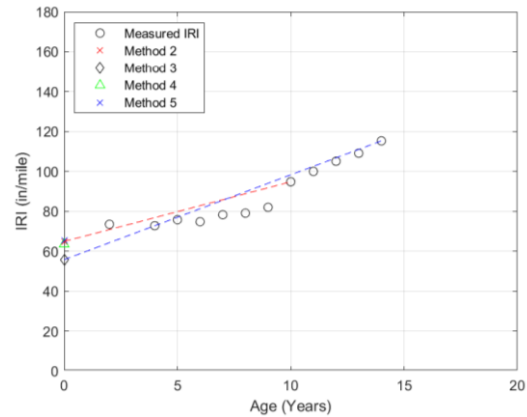
(a) Section 1



(b) Section 2



(c) Section 3



(d) Section 4

Figure 3-27 Illustration showing backcasting of initial IRI

Table 3-16 Recommended thresholds based on initial IRI for flexible sections

IRI less than or equal to	No of sections	Mean initial IRI (in/mile)
85	380	56.1
82	371	55.4
80	362	54.8
78	356	54.4
75	349	53.9
70	331	52.9
67	295	51.0
65	274	49.7

Table 3-17 Recommended thresholds based on initial IRI for rigid sections

IRI less than or equal to	No. of sections	Mean initial IRI (in/mile)
85	74	73.7
82	65	71.6
80	52	69

Table 3-18 Summary of initial IRI thresholds

Pavement type	Fix type	Initial IRI threshold (in/mile)
Flexible	New	77
	Overlay	82
Rigid	New	82
	Overlay	82

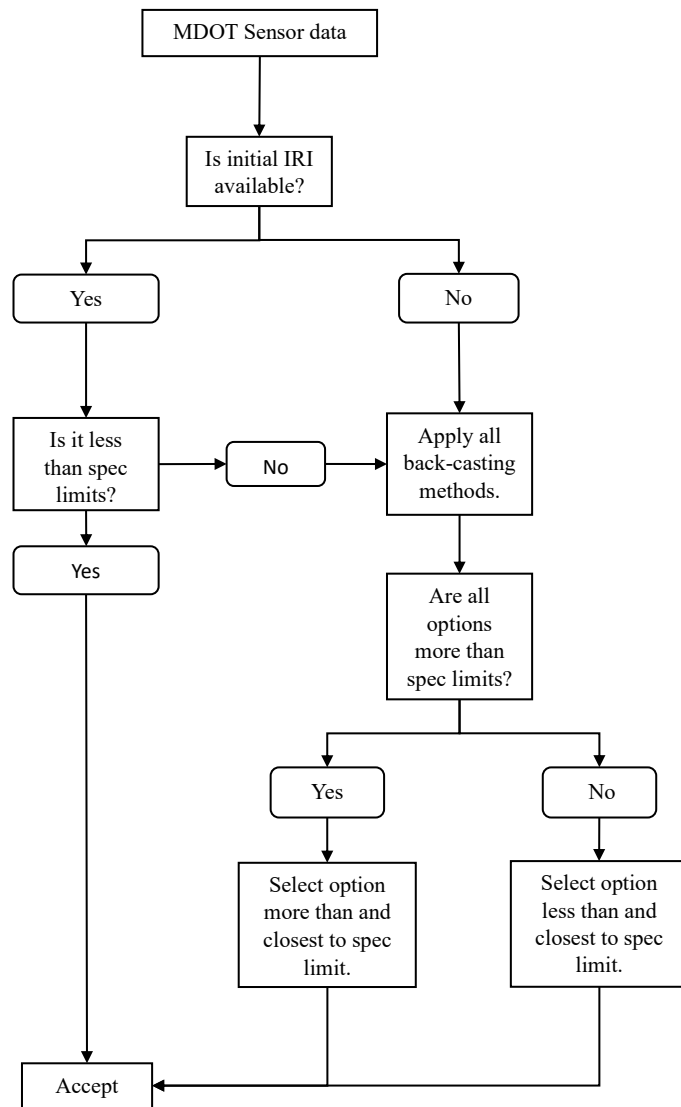


Figure 3-28 Flowchart for selection of initial IRI

Therefore, some thresholds were selected to keep reasonable initial IRI values. Any section with an initial IRI value higher than the threshold was eliminated from the IRI calibration. Tables 3-16 and 3-17 show different threshold values for flexible and rigid sections. It is important to note that the sections in Tables 3-16 and 3-17 consist of both new and overlay sections, but only new

sections have been used in this study. Based on the number of sections available and the average IRI for each cap, different threshold limits were selected for flexible and rigid pavements, as shown in Table 3-18. Figure 3-29 shows the distribution of initial IRI for the selected flexible and rigid sections. The distribution of the initial IRI is acceptable for an optimum IRI model calibration.

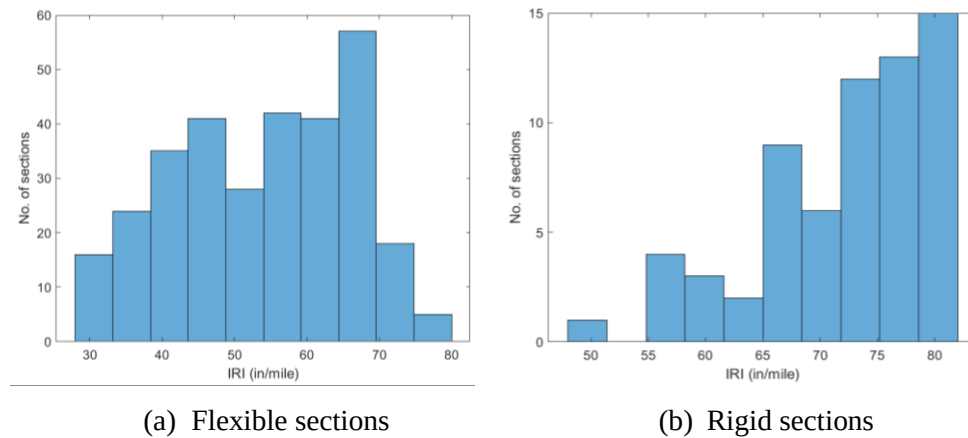


Figure 3-29 Distribution of initial IRI

It is essential to verify the accuracy of the proposed methodology. Five sections from flexible and one from rigid are taken for this purpose. These sections have the initial IRI data available at zero (construction) year. Only one rigid section has initial IRI data available at zero year. Methods 2 to 5 are implemented using measured IRI data from age 1 to 20 (excluding zero-year data). The comparison between the recommended initial IRI based on the proposed methodology and the recorded initial IRI shows a good correlation with an error of less than 8% for all sections. Table 3-19 shows the summary of the validation results.

Table 3-19 Summary of validation results

Pavement type	Initial IRI backcasting (in/mile)				Recommended Initial IRI (in/mile)	Recorded initial IRI (in/mile)	Error (%)
	Method 2	Method 3	Method 3	Method 4			
Flexible	40.6	40.6	32.9	36.9	40.6	41.2	1.4
	57.7	57.7	52.7	56.7	57.7	55.5	4.0
	42.6	42.6	40.5	44.5	44.5	45.7	2.6
	35.8	34.2	28.9	32.9	35.8	38.9	8.0
	55.5	50.9	48.4	52.4	55.5	58.4	5.0
Rigid	70.7	66.6	52.0	56.0	70.7	72.4	2.3

3.6 CHAPTER SUMMARY

This chapter outlines the data used for local calibration, emphasizing the importance of selecting representative pavement sections and gathering pertinent data for accurate performance predictions. It details the methodology of converting the MDOT PMS data to Pavement-ME compatible units, evaluating distress trends, and considering maintenance history. Key distresses were identified, and databases were created for efficient data extraction. Project selection criteria prioritize sections with adequate data and performance trends. The selected sections were also verified against all MDOT sections to validate if these sections are representative of overall MDOT performance. Sections were categorized as good, fair, or poor based on measured trends relative to reference lines. The results showed that the selected sections represent MDOT pavement sections well. A total of 256 flexible and 88 rigid sections were selected. The number of projects for each performance type and pavement type has also been summarized. This chapter also details each input, source, and possible estimates in case of unavailable data. These inputs include the HMA and PCC material inputs, traffic, climate, and estimation of initial IRI. Table 3-20 summarizes the inputs and corresponding levels for traffic, climate, and material characterization data used for the local calibration.

Table 3-20 Summary of input levels and data source

Input			Pavement-ME input level	Data source level	Input source
Traffic	Vehicle class distribution		1	2	MDOT specified traffic per cluster data
	Hourly distribution		1	2	
	Monthly adjustment factor		1	2	
	Number of axles per truck		1	2	
	Single, tandem, tridem, and quad axle load distribution		1	2	
	AADTT		1	1	From design drawings
	Vehicle class 9 percentage		1	1	MDOT TDMS website
Cross-section layers (new and existing)	HMA thickness		1	1	Project-specific HMA thicknesses based on design drawings
	PCC thickness		1	1	Project-specific PCC thicknesses based on design drawings
	Base thickness		1	1	Project specific base thicknesses based on design drawings
	Subbase thickness		1	1	Project-specific subbase thicknesses based on design drawings
Layer materials	HMA	Mix properties	1	Mix of 2 and 3	MDOT HMA mixture characterization study (DYNAMOD database)
		HMA mixture aggregate gradation	1	1 or 3	Project-specific mixture gradation data obtained from data collection or average statewide values
		Binder properties	1	3	MDOT HMA mixture characterization study (DYNAMOD database)
	PCC	Strength (f'_c , MOR)	3	1 or 3	Project specific testing values or average statewide value
		CTE	1	2	MDOT recommended values
	Base/subbase	MR	3	3	Recommendations from MDOT unbound material study
	Subgrade	MR	3	3	Soil-specific MR values per MDOT subgrade soil study
		Soil properties	Mix of all levels	3	Location-based soil type per MDOT subgrade soil study
Climate			1	1	Closest available climate station

Note:

Data source Level 1 is project-specific data

Data source Level 2 inputs are based on regional averages in Michigan

Data source Level 3 inputs are based on statewide averages in Michigan

CHAPTER 4 - METHODOLOGY

4.1 INTRODUCTION

Local calibration of the Pavement-ME models aims to optimize the model coefficients by minimizing bias and standard error, which is achieved by matching the predicted and measured distress. Bias in the predictions signifies if there is a systematic over or under-prediction, whereas standard error shows the scatter and variability. Figure 4-1 shows a representation of bias and standard error. This chapter highlights each model's calibration methods and approaches, the reliability calculation, and the sensitivity analysis of Pavement-ME model coefficients.

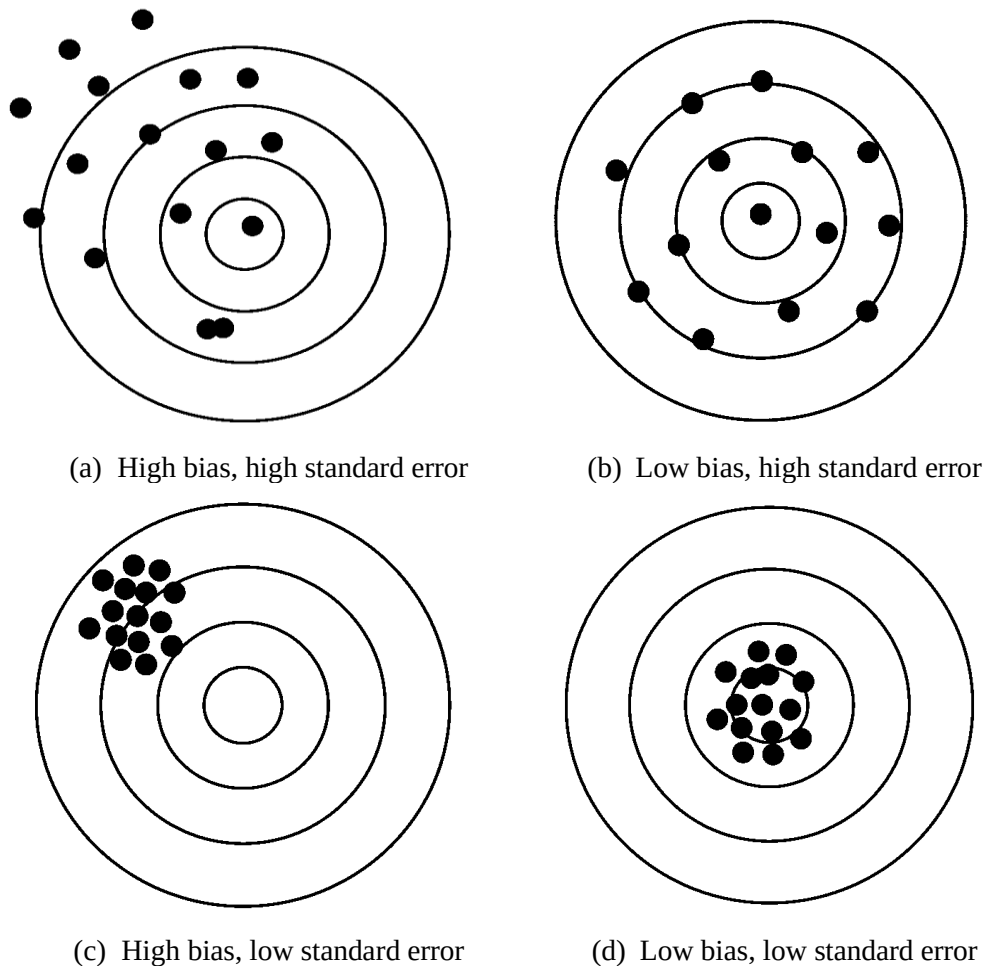


Figure 4-1 Schematic representation of bias and standard error (10)

The details for inputs, performance data, and project selection have already been discussed in Chapter 3. Once the data is extracted, it can be used to run the Pavement-ME files (.dgp files) and generate outputs (structural responses). The process for local calibration is summarized below:

- (a) Run the Pavement-ME (using global model coefficients) and extract critical responses and predicted distresses.
- (b) Compare the predicted distress with measured distress.
- (c) Based on the results from Step 2, test the accuracy of the global models and the need for local calibration.
- (d) If predictions using global models are satisfactory, local calibration is not required, and global models can be accepted. Local calibration is needed if the global model has significant bias and standard error.
- (e) Check your calibration results by validating them on an independent set of sections not used for calibration.
- (f) Estimate the reliability equations based on the calibrated model predictions and measured distress.

Before locally calibrating the Pavement-ME models, it is vital to determine the need for calibration. This includes testing the accuracy of the global model predictions at a reliability of 50%, which is the mean prediction. Once the predictions from the global model are obtained, they are compared with measured values to calculate bias and standard error. A plot of predicted versus measured values is created for each distress to visualize the accuracy of predictions to a line of equality (LOE). Testing the global model also includes hypothesis testing. For a good fit, the points should lie along the LOE. The measured distress $y_{Measured}$ and predicted distress $x_{Predicted}$ can be modeled as a linear model as shown in Equation (4-1), where m is the slope, and b_o is the intercept.

$$y_{Measured} = b_o + m \times x_{Predicted} \quad (4-1)$$

Three hypothesis tests are conducted to evaluate the reasonableness of the global model. If any of these hypotheses fail, the models are recalibrated for local conditions:

- There is no systematic bias between the measured and predicted distress [Equation (4-2)]. This can be tested using a paired t-test.

$$H_0: \sum (y_{\text{Measured}} - x_{\text{Predicted}}) = 0 \quad (4-2)$$

- The slope parameter m is 1 (Equation (4-3)).

$$H_0: m = 1.0 \quad (4-3)$$

- The intercept parameter b_o is zero (Equation (4-4)).

$$H_0: b_o = 0 \quad (4-4)$$

4.2 CALIBRATION APPROACHES

The empirical Pavement-ME transfer functions can be of two types: (a) model that directly calculates the magnitude of surface distress, and (b) model that calculates the cumulative damage index rather than actual distress magnitude.

Approach 1: For specific models (e.g., fatigue cracking, rutting, transverse cracking, and IRI), damage is directly obtained from Pavement-ME outputs. The transfer functions predict distress from the damage and have been calibrated using the MATLAB program outside the Pavement-ME. Different resampling techniques and MLE have been used to calibrate these functions. Genetic Algorithm (GA) has been used to optimize transfer function coefficients using MATLAB program for this approach. These MATLAB codes are available from the author upon request. GA is an evolutionary optimization technique that can converge towards a global minimum solution even with local minima. GA involves the following operations:

- Initialization: GA generates solutions by randomly selecting a subset inside the allowed search space called the population.
- Selection: The generated solutions are selected based on the value of the objective function.
- Generation of offspring: New solutions are created using the selected solutions or populations (offspring) based on two main processes: mutation and crossover.
- Termination: This process continues till the termination criteria for the given population or the number of generations is reached.

Approach 2: The Calibration Assistance Tool (CAT) calibrated the models (e.g., thermal cracking and joint faulting) where the damage is not obtained from the Pavement-ME outputs. These models predict distress by calculating cumulative damage over time. One can't use the resampling techniques or the MLE method for this approach.

Based on the model, two different calibration approaches have been followed (as shown in Table 4-1):

Table 4-1 Model transfer functions and calibration approaches (28)

Pavement type	Performance prediction model	Approach		Model transfer functions
		I	II	
Flexible pavement	Fatigue cracking – bottom up	✓		$FC_{Bottom} = \left(\frac{1}{60}\right) \left(\frac{6000}{1 + e^{C_1 C_1^* + C_2 C_2^* \log(DI_{Bottom} \cdot 100)}} \right)$
	Fatigue cracking – top down	✓		$t_0 = \frac{K_{L1}}{1 + e^{K_{L2} \times 100 \times (a_0/2A_0) + K_{L3} \times HT + K_{L4} \times LT + K_{L5} \times \log_{10} AADTT}}$ $L(t) = L_{MAX} e^{-\left(\frac{C_1 \rho}{t - C_3 t_0}\right)^{C_2 \beta}}$
	Rutting	HMA	✓	$\Delta_{p(HMA)} = \varepsilon_{p(HMA)} h_{HMA} = \beta_{1r} k_z \varepsilon_r(HMA) 10^{k_{1r}} n^{k_{2r}} \beta_{2r} T^{k_{3r}} \beta_{3r}$
		Base/subgrade	✓	$\Delta_{p(soil)} = \beta_{s1} k_{s1} \varepsilon_v h_{soil} \left(\frac{\varepsilon_0}{\varepsilon_r}\right) e^{-\left(\frac{\rho}{n}\right)^\beta}$
	Thermal cracking		✓	$A = 10^{k_t \beta_t (4.389 - 2.52 \log(E_{HMA} \sigma_m \eta))}$
	IRI	✓		$IRI = IRI_o + C_1(RD) + C_2(FC_{Total}) + C_3(TC) + C_4(SF)$
Rigid pavement	Transverse cracking	✓		$CRK_{BU/TD} = \frac{100}{1 + C_4(DI_F)^{C_5}}$ $TCRACK = (CRK_{Bottom-up} + CRK_{Top-down} - CRK_{Bottom-up} \cdot CRK_{Top-down}) \cdot 100\%$
	Transverse joint faulting		✓	$Fault_m = \sum_{i=1}^m \Delta Fault_i$ $\Delta Fault_i = C_{34} \times (FAULTMAX_{i-1} - Fault_{i-1})^2 \times DE_i$ $FAULTMAX_i = FAULTMAX_0 + C_7 \times \sum_{j=1}^m DE_j \times \log(1 + C_5 \times 5.0^{EROD})^{C_6}$ $FAULTMAX_0 = C_{12} \times \delta_{curling} \times \left[\log(1 + C_5 \times 5.0^{EROD}) \times \log\left(\frac{P_{200} \times WetDays}{p_s}\right) \right]^{C_6}$ $C_{12} = C_1 + C_2 \times FR^{0.25}$ $C_{34} = C_3 + C_4 \times FR^{0.25}$
	IRI	✓		$IRI = IRI_o + C_1(CRK) + C_2(SPALL) + C_3(TFAULT) + C_4(SF)$

*Bold font indicates calibration coefficients

4.3 CALIBRATION METHODS

This study used different methods (least squares and MLE) to demonstrate and compare calibration differences for normally and non-normally distributed data. For example, the measured transverse cracking in rigid pavements is typically non-normally distributed, with most data points near zero, whereas IRI is close to a normal distribution. Both methods have their advantages and limitations. It is important to note that thermal cracking, top-down cracking in flexible pavements, and joint faulting in rigid pavements were not calibrated using the MLE method.

The measured data is limited to the MDOT PMS database. Apart from the measured data, this study also used synthetic data as it provides the freedom to generate any distribution with random errors. This methodology also validates a more generic use of MLE on a dataset outside the measured data. Before calibration using measured data, synthetic data was created to show the applicability of the MLE approach. For this purpose, DI_{Bottom} was generated using an exponential distribution with $\lambda = 0.3$ to generate synthetic bottom-up cracking data in flexible pavements. DI_{Bottom} was used to calculate bottom-up cracking for 355 points, the same number of points as the measured data. A value of $C_1 = 0.254$, $C_2 = 0.730$ (for total AC thickness (T) < 5 in.), and $C_2 = (0.867 + 0.2583 * T) * 0.238$ (for 5 in. ≤ T ≤ 12 in.) were used for calculation of bottom-up cracking. The assumption of an exponential distribution and the value of λ is based on the measured bottom-up cracking data. The generated synthetic data is close to the measured data but follows a smooth exponential distribution curve. Two different datasets were created, one without variability (no change introduced in the generated data) and one with a uniformly distributed random variability applied on each data point between -50 % and 50%. A similar methodology created synthetic data for transverse cracking in rigid pavements. Initially, an exponentially distributed DI_F was generated using $\lambda = 0.1$. The generated DI_F was then used to calculate transverse cracking. About 237 points were generated for the synthetic data, the same as for measured transverse cracking data.

The selection of a suitable method and distribution is based on several parameters. Negative log-likelihood (NLL) was calculated for the MLE and least squares methods, the formulation for which is presented in the proceeding sections. Besides the NLL values, four other statistical parameters were used as selection criteria for the most suitable model. These are the Standard Error of Estimate (SEE), bias, Akaike Information Criterion (AIC), and Bayesian Information Criterion (BIC). AIC is a statistical measure used for model selection that balances the goodness of fit with the complexity of the model, as shown in Equation (4-5). BIC is a similar criterion that penalizes model complexity more strongly, often leading to more efficient model selection, as shown in Equation (4-6).

$$AIC = 2S - 2LL \quad (4-5)$$

$$BIC = \ln(n)S - 2LL \quad (4-6)$$

where,

n = Number of data points

S = Number of parameters of distribution (for example $S = 1$ for exponential distribution)

LL = Log-likelihood value

4.3.1 Calibration Using the Least Squares Method

The least squares method is a popular technique used in various statistics, mathematics, and engineering fields to fit mathematical models to data. Its primary aim is to minimize the sum of the squares of the residuals between observed and predicted values. It follows the NIID assumption, which may not apply to non-normally distributed data. This method was employed to estimate the parameters of the Pavement-ME transfer functions. The fundamental idea behind the least squares method is to find the line (or curve) that best fits a set of data points by minimizing the sum of the squared differences between the observed data points and the corresponding values predicted by the model. The bias and SEE values were minimized using the least squares method, as shown in Equations (4-7) and (4-8)

$$SEE = \sqrt{\frac{\sum (y - \hat{y})^2}{n - 1}} \quad (4-7)$$

$$Bias = \sum (y - \hat{y}) \quad (4-8)$$

where,

y = Measured data

$\hat{y}$ = Predicted data

n = Number of data points

4.3.2 Calibration Using the Maximum Likelihood Estimation (MLE) Method

MLE is a powerful statistical technique for parameter estimation in various fields, including biology, physics, economics, and engineering. MLE was used to calibrate the bottom-up cracking, total rutting, and IRI models in flexible pavement and transverse cracking and IRI models in rigid pavements. MLE seeks to estimate the parameters of a probability distribution that best describes the observed data based on the likelihood function. The likelihood function measures the probability of the observed data following a known distribution. MLE finds the set of distribution parameters that maximize the likelihood function. Consider a dataset $X = (x_1, x_2,$

..., x_n) that is generated by a probability distribution with parameters θ . The likelihood function $L(\theta|X)$ is the joint probability density function of the observed data, given the distribution parameters as shown in Equation (4-9).

$$L(\theta|X) = P(X|\theta) = P(x_1, x_2, \dots, x_n|\theta) \quad (4-9)$$

Here, P denotes the probability density function, and the likelihood function measures the probability of observing the data X given the distribution parameters θ . The goal of MLE is to find the set of distribution parameters θ that maximizes the likelihood function between dataset X and the assumed distribution. In practice, it is often easier to work with the log-likelihood function so that the product of likelihood values becomes a summation; one can do this by taking the natural logarithm of the likelihood function. The log-likelihood function is given by Equation (4-10).

$$l(\theta|X) = \log L(\theta|X) = \log P(X|\theta) = \log \prod P(x_i|\theta) = \sum \log P(x_i|\theta) \quad (4-10)$$

where;

$\prod$ = Product operator

$\sum$ = Summation operator

Taking the logarithm of the likelihood function also simplifies the computation of the derivative required for optimization. One can solve the optimization problem by finding the values of θ that maximize the log-likelihood function. This can be done using numerical optimization algorithms, such as gradient descent, Newton's, or quasi-Newton methods. These algorithms require the derivative of the log-likelihood function for the distribution parameters. Numerical optimization algorithms iteratively update the values of the distribution parameters to find the maximum of the log-likelihood function. The optimization process continues until the algorithm converges to a maximum of the log-likelihood function. The MLE obtained from the optimization process represents the most likely estimates of the distribution parameters that can explain the observed data. The calibration process for MLE involves the following steps:

Step 1: Assume the initial values of the transfer function coefficients to calculate the predicted cracking.

Step 2: Fit a known distribution (for example, exponential, gamma, etc.) to the predicted cracking and estimate the distribution parameters.

Step 3: Calculate the NLL between the known distribution parameters in Step 2 and the measured values.

Step 4: Repeat Steps 1 to 3 to minimize the NLL value.

Step 5: Coefficients with minimum NLL are the desired coefficients.

Four distributions were used for this analysis: gamma, log-normal, exponential, and negative binomial. The Probability Density Function (pdf)/ Probability Mass Function (pmf) of these distributions is shown in Equation (4-11) to (4-14), respectively.

- Gamma distribution

$$f(x) = \frac{x^{\alpha-1} e^{-x/\beta}}{\beta^\alpha \Gamma(\alpha)} \quad (4-11)$$

- Log-normal distribution

$$f(x) = \frac{e^{-\left(\frac{(\ln((x-\theta)(m))^2/(2\sigma^2))}{(x-\theta)\sigma\sqrt{2\pi}}\right)}}{(x-\theta)\sigma\sqrt{2\pi}} \quad x > \theta; m, \sigma > 0 \quad (4-12)$$

- Exponential distribution

$$f(x) = \lambda e^{-\lambda x} \quad (4-13)$$

- Negative binomial distribution

$$P(X = x | r, p) = \binom{x-1}{r-1} p^r (1-p)^{x-r}, \quad x = r, r+1, \dots, \quad (4-14)$$

The formulation of the maximum likelihood function for exponential distribution is shown below. A similar approach was used for other distributions. Equation (4-15) shows the pdf for exponential distribution. Comparing it with Equation (4-13), here $\lambda = \frac{1}{\beta}$, which is the rate parameter, and x is the observed value. The likelihood function for a set of independent and identically distributed observations from the exponential distribution is obtained by taking the product of the individual probability density functions shown in Equations (4-16) and (4-17).

$$f(\mathbf{x}, \beta) = \frac{1}{\beta} e^{\left(\frac{-x}{\beta}\right)}; \mathbf{x} > 0 \quad (4-15)$$

$$L(\beta, \mathbf{x}) = L(\beta, x_1, \dots, x_N) = \prod_{i=1}^N f(x_i, \beta) \quad (4-16)$$

$$L(\beta, \mathbf{x}) = \prod_{i=1}^N \frac{1}{\beta} e^{\left(\frac{-x_i}{\beta}\right)} \quad (4-17)$$

It is common to work with the log-likelihood function instead of the likelihood function to simplify the calculation. The log-likelihood function is obtained by taking the natural logarithm of the likelihood function, as shown in Equation (4-18).

$$\mathcal{L}(\beta, \mathbf{x}) = \log \left(\prod_{i=1}^N \frac{1}{\beta} e^{\left(\frac{-x_i}{\beta}\right)} \right) \quad (4-18)$$

Simplifying Equation (4-18) using properties of the log is shown in Equations (4-19) to (4-21). Equation (4-21) shows the negative log-likelihood of exponential distribution used for calibration.

$$\mathcal{L}(\beta, \mathbf{x}) = \log \left(\prod_{i=1}^N \frac{1}{\beta} e^{\left(\frac{-x_i}{\beta}\right)} \right) = \sum_{i=1}^N \left(\log \left(\frac{1}{\beta} \right) + \log \left(e^{\left(\frac{-x_i}{\beta}\right)} \right) \right) \quad (4-19)$$

$$\mathcal{L}(\beta, \mathbf{x}) = N \log \left(\frac{1}{\beta} \right) + \sum_{i=1}^N \left(\frac{-x_i}{\beta} \right) \quad (4-20)$$

$$\mathcal{L}(\beta, \mathbf{x}) = -N \log (\beta) + \frac{1}{\beta} \sum_{i=1}^N -x_i \quad (4-21)$$

To estimate the value of β at the maxima of log-likelihood, Equation (4-21) can be differentiated. Equations (4-22) to (4-24) show the estimation of β at the maxima of log-likelihood.

$$\frac{\partial \mathcal{L}}{\partial \beta} = \frac{\partial}{\partial \beta} \left(-N \log (\beta) + \frac{1}{\beta} \sum_{i=1}^N -x_i \right) = 0 \quad (4-22)$$

$$\frac{\partial \mathcal{L}}{\partial \beta} = -\frac{N}{\beta} + \frac{1}{\beta^2} \sum_{i=1}^N x_i = 0 \quad (4-23)$$

$$\beta = \frac{\sum_{i=1}^N x_i}{N} = \bar{\mathbf{x}} \quad (4-24)$$

Figure 4-2 shows the flow chart of the methodology used and the final selection of the optimum method and distribution.

4.4 RESAMPLING TECHNIQUES

Various sampling techniques were used to calibrate Pavement-ME transfer functions. The least squares and MLE methods were combined with these techniques to improve the robustness of

the estimated parameters. All these techniques have been used for models calibrated using Approach 1. For models calibrated using Approach 2, no sampling or traditional split sampling has been used in the CAT tool.

1. *No sampling*: This technique considers the entire dataset (all available measured data points and corresponding damage) and was used for both Approaches 1 and 2.
2. *Traditional split sampling*: The dataset is randomly divided into two parts—70% of the data for the calibration set and the rest 30% for the validation set. The optimization is performed only on the calibration set, and the obtained coefficients are applied to an independent validation set. This method was used for both Approaches 1 and 2.
3. *Repeated split sampling*: This technique is like traditional split sampling but with 1000 resamples, where a different data set was picked up each time for calibration (70%) and validation (30%). This method was used only for Approach 1.
4. *Bootstrapping*: Bootstrap resampling is used to draw 1000 bootstrap samples from the original dataset with replacement. Each bootstrap resamples the original data with the same sample size but may contain some duplicate observations. This method estimates a sampling distribution and confidence intervals for a population parameter, even when the underlying population distribution is unknown.

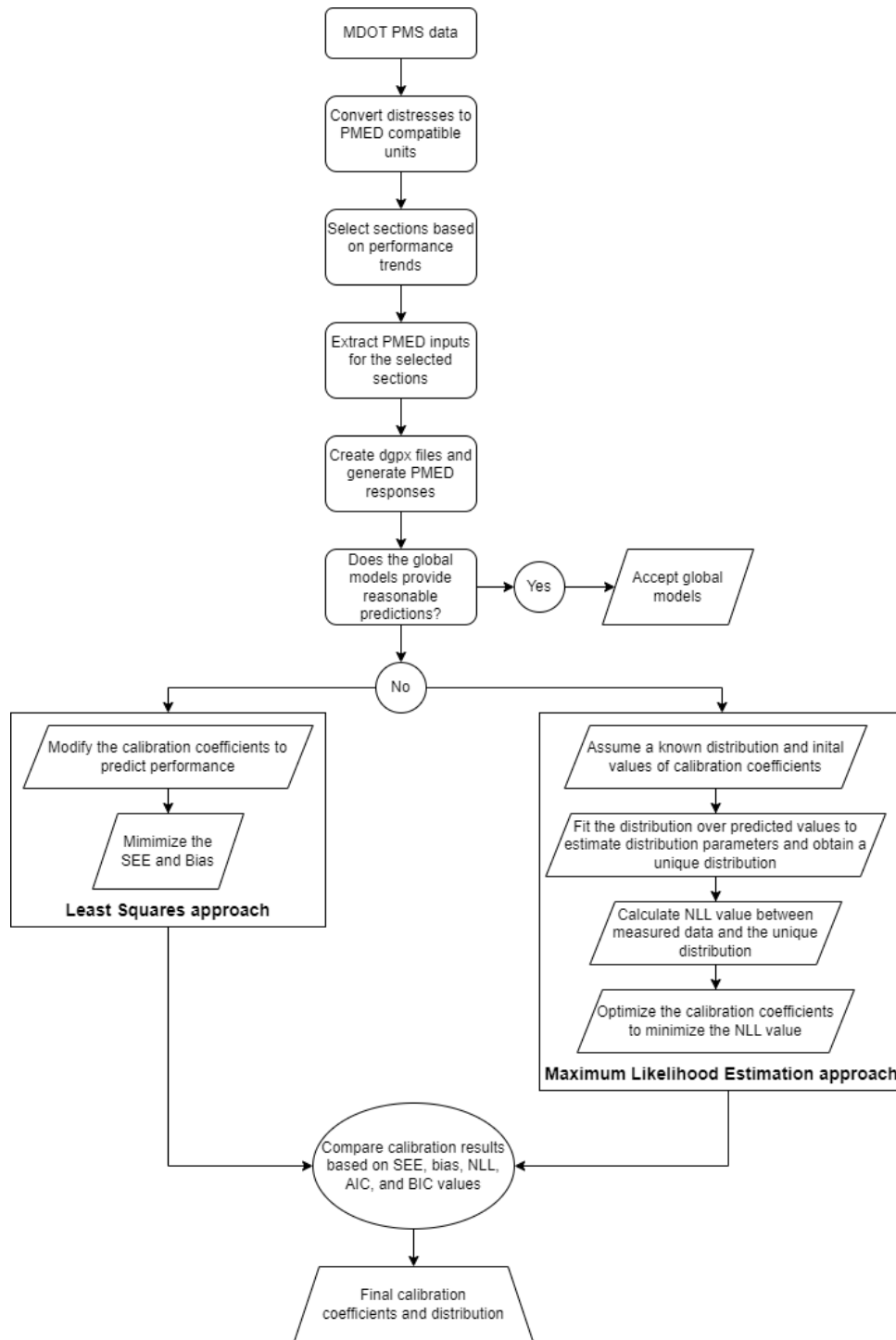


Figure 4-2 Flowchart of calibration methodology

Traditional no-sampling or split sampling technique provides a convenient approach to selecting pavement sections from the calibration database. Though these techniques are easy to implement and can be used for any Pavement-ME model, they might impose some limitations. Resampling

techniques have several advantages over traditional approaches. Since these are non-parametric techniques, the model parameters can be estimated without making assumptions about the data distribution. The distribution of the model coefficients and error parameters can be estimated instead of the point estimate. This can give a better estimation of parameters within desired confidence intervals. Since a new sample is created every time, the outliers or sections controlling the calibration process can be identified. Though these resampling techniques have several advantages over traditional approaches, there are also certain limitations. Bootstrapping cannot be used for small datasets or when the independence assumption is unmet. Resampling techniques also require higher computing power and time and can be used only for those performance models where the damage and other inputs are available from Pavement-ME. Table 4-2 summarizes the advantages and limitations of all calibration techniques.

Table 4-2 Summary of calibration techniques

Technique	Advantages	Limitations
No sampling	• Computationally efficient • Applicable even for small sample size	• Provides point estimates • It may not be suitable for non-normally distributed data
Split sampling	• Computationally efficient • Provides validation	• Provides point estimates • It may not be suitable for non-normally distributed data
Repeated split sampling	• Provides confidence intervals • Provides validation • Identifies outliers • Distribution assumption is not required	• Computationally time-consuming • It cannot be used for smaller sample size • It may not be suitable for non-normally distributed data
Bootstrapping	• Provides confidence intervals • Identifies outliers • Distribution assumption is not required	• Computationally time-consuming • It cannot be used for smaller sample size • It may not be suitable for non-normally distributed data

4.5 FLEXIBLE PAVEMENT MODEL COEFFICIENTS

The design distress in the Pavement-ME includes bottom-up cracking, top-down cracking, rutting, thermal (transverse) cracking, reflective cracking, and IRI. The calibration of each model and the specific coefficients calibrated has been discussed in this section.

4.5.1 Fatigue Cracking Model (Bottom-up)

The fatigue cracking (bottom-up) model was calibrated by optimizing the C_1 and C_2 coefficients (see Table 4-1). In Pavement-ME v2.6, coefficient C_1 is a single value, whereas coefficient C_2 has three different values depending on the total HMA thickness. Table 4-3 shows the global values for C_1 and C_2 .

Table 4-3 Global values for bottom-up cracking model coefficients

Calibration coefficient	Global values
C_1	1.31
C_2	$H_{ac} < 5 \text{ in.} : 2.1585$
	$5 \text{ in.} \leq H_{ac} \leq 12 \text{ in.} : (0.867 + 0.2583 \times H_{ac}) \times 1$
	$H_{ac} > 12 \text{ in.} : 3.9666$

H_{ac} : Total HMA thickness in inches

Notably, no sections were selected for the bottom-up calibration with a total HMA thickness of more than 12 inches. The coefficient C_2 was calibrated separately for the thickness ranges less than 5 inches and 5 to 12 inches, respectively. For a thickness range of 5 to 12 inches, only the multiplying factor 1 (marked in bold here: $(0.867 + 0.2583 \times H_{ac}) \times \mathbf{1}$) was calibrated, while other values (0.867 and 0.2583) were kept at global values. A single value was used for a thickness range of more than 12 inches. The H_{ac} was kept at 12 inches, and the multiplying factor 1 was kept at the calibrated value obtained for the 5 to 12-inch thickness range. The crack initiation time is affected by C_1 , whereas the slope of the bottom-up cracking curve is affected by C_2 . Consequently, the calibration was performed using two approaches: (a) combined measured bottom-up and top-down cracking and (b) bottom-up cracking only. MLE was used for approach (a), whereas least squares was used for both methods.

4.5.2 Fatigue Cracking Model (Top-down)

The top-down cracking model has been modified in the Pavement-ME v2.6. The model consists of a crack initiation function that calculates the time to crack initiation and a crack propagation function that calculates the percent lane area cracked. This makes it a total of eight coefficients combined from both functions. Since the actual crack initiation time was not known, it was not possible to calibrate the crack initiation model separately. So, a single function was used by substituting the crack initiation function with the crack propagation function. Initially, an attempt was made to change all eight coefficients simultaneously. This approach had some challenges:

- The model has some mathematical limitations. High values for C_3 cause mathematical errors when using it in Pavement-ME.
- No current literature exists for the top-down cracking model calibration. Therefore, estimating the range for each coefficient to be used in optimization was difficult.
- The model has many coefficients with coefficient values ranging from 0.011 to 64271618. This makes the optimization challenging to converge.

As mentioned above, four coefficients from the crack initiation function (kL2, kL3, kL4, kL5) and two coefficients from the crack propagation function (C1, C2) have been calibrated based on the model's understanding and limitations.

4.5.3 Rutting Model

Due to axle loads, rutting is the total accumulated plastic strain in different pavement layers (HMA, base/sub-base, and subgrade). It is calculated by summing up the plastic strains at the mid-depth of individual layers accumulated for each time increment. In the Pavement-ME, rutting is predicted separately for the layers (HMA, base, and subgrade). The total rutting is the sum of rutting from all layers. The AC rutting model has three coefficients (β_{1r} , β_{2r} , β_{3r}). β_{1r} is a direct multiplier and was calibrated using optimization outside the Pavement-ME. In this model, β_{2r} and β_{3r} are power to the pavement temperature and the number of axle load repetitions. Calibration of β_{2r} and β_{3r} cannot be done outside of the Pavement-ME and requires running the Pavement-ME multiple times or optimizing these in the CAT tool. Initially, β_{2r} and β_{3r} values were used from the previous calibration effort, and β_{1r} was calibrated (10). This calibration approach provided reasonable results; therefore, β_{2r} and β_{3r} from the previous calibration were accepted, and only β_{1r} was calibrated.

The unbound layers (base and subgrade) rutting model have one calibration coefficient each (β_{s1}). Since β_{s1} is a direct multiplier, it can be calibrated using optimization outside the Pavement-ME without running the software or CAT tool. Since both base and subgrade have the same model and calibration coefficient, the base calibration coefficient is referred to as β_{s1} , and the subgrade coefficient is referred to as β_{sg1} . The rutting model in the Pavement-ME was calibrated using the following two methods:

- *Method 1: Individual layer rutting calibrations* — The measured rutting from individual layers was matched against the Pavement-ME predictions (β_{1r} , β_{s1} , and β_{sg1} were

calibrated separately) for this approach. The total measured rutting was multiplied by the percent contribution from each layer to obtain measured rutting for the individual layer. Figure 4-3 shows the percentage contribution estimated using transverse pavement profile analysis. The width and depth of the measured rut channel were used to determine the seat of rutting and rutting in individual layers. AC layer rutting contributes more than 70% to all pavement types [based on transverse profile analysis (10)]. Pavement-ME has separate standard error equations for rutting in the individual layers. This method evaluated the standard error equations for rutting in each layer.

- **Method 2: Total surface rutting calibration** — The total measured rutting was calibrated against the sum of individual predicted rutting (i.e., β_{1r} , β_{s1} , and β_{sg1} were calibrated simultaneously).

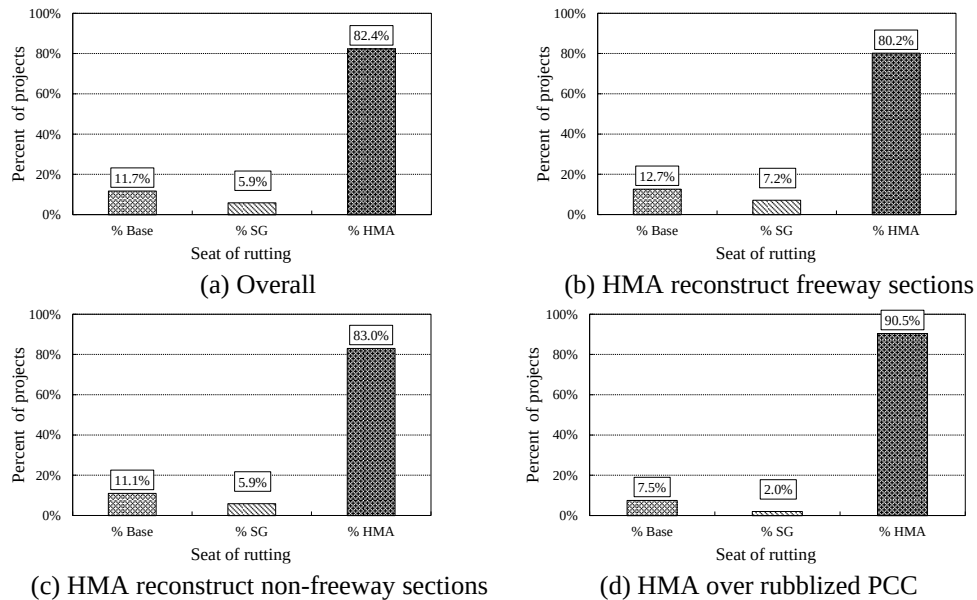


Figure 4-3 Transverse profile analysis for total rutting (10)

4.5.4 Thermal Cracking Model

The thermal cracking model in the Pavement-ME has three different levels for the calibration coefficient. These levels are based on the level of HMA input. Level 1 G^* and Level 2 IDT have been used to calibrate the thermal cracking model. This corresponds to Level 1 thermal cracking calibration coefficients. Both G^* and IDT values were obtained from the DYNAMOD software database. In the DYNAMOD database, G^* and IDT values are available only for sections with Performance grade (PG) binder type. Therefore, sections with PG binder type (Superpave mixes)

have been used to calibrate the thermal cracking model. In the Pavement-ME v2.6, the calibration coefficient k_t is originally a function of the mean annual air temperature (MAAT), whereas, in v2.3, it was a single representative value. Two different approaches were used for calibration:

- (a) Using the CAT tool, an initial attempt was made to calibrate k_t (using the original equation as a function of MATT).
- (b) A second attempt was made to calibrate k_t by running the Pavement-ME multiple times with different k_t values of 0.25, 0.65, 0.75, 0.85, 0.95, and 1.35. This time, single values for k_t were used, which were not a function of MAAT.

k_t as a function of MAAT resulted in contradictory results when comparing Michigan temperature extremes, where thermal cracking at cold temperatures was either reduced or equal to thermal cracking at warm temperatures. Moreover, ARA recommends using a single k_t value if this is more suitable for the agency and its local conditions. Based on these results, the k_t value based on the second approach was recommended. It is important to note that for this calibration, the average thermal cracking for a section was cut at 2112 ft/mile.

4.5.5 IRI Model for Flexible Pavements

IRI is a linear function of initial IRI, rut depth, total fatigue cracking, transverse cracking, and site factor. The initial IRI was obtained from linear backcasting based on the time series trend for each section, as described in Chapter 3. The fatigue cracking, rutting, and transverse cracking models were calibrated before calibrating the IRI model. Since all inputs to the IRI model could be obtained, it was calibrated outside Pavement-ME. IRI has a closed-form solution and does not require a standard error equation in the Pavement-ME. The standard error for IRI is calculated using the standard error of its components.

4.6 RIGID PAVEMENT MODEL COEFFICIENTS

The design distresses in the Pavement-ME include transverse cracking (percentage of slabs cracked), transverse joint faulting (inches), and international roughness index (IRI) for rigid pavements. The calibration methodology for each model is discussed in this section.

4.6.1 Transverse Cracking Model

The coefficients C_4 and C_5 (shown in Table 4-1) were optimized to calibrate the transverse cracking model. These coefficients were calibrated outside the Pavement-ME and without the CAT tool. C_4 affects the crack initiation time, and C_5 affects the slope of the transverse cracking curve.

4.6.2 Transverse Joint Faulting Model

The joint faulting model in the Pavement-ME consists of eight coefficients. Joint faulting could not be predicted using the available inputs outside the Pavement-ME; therefore, it was calibrated using the CAT tool. CAT tool has a limitation on the run time and the total combinations of coefficients that can be calibrated simultaneously. Therefore, it was essential to identify the most sensitive coefficients. Several research studies (11, 26) show that out of the eight calibration coefficients for the faulting model, C_6 is the most sensitive. C_1 is the next sensitive coefficient, followed by C_2 . Using this sequence of sensitivity of the different coefficients, C_1 and C_6 were calibrated together. The calibrated coefficients from C_1 and C_6 were kept fixed, and C_2 was calibrated. In this sequence, the three most sensitive coefficients were calibrated. As previously noted and explained in Chapter 3, the joint faulting (for every 0.1-mile segment) was cut at 0.4 inches for calibration.

4.6.3 IRI Model for Rigid Pavements

IRI in rigid pavements is a linear function of initial IRI, transverse cracking, joint spalling, faulting, and site factor. The initial IRI was obtained from linear backcalculation based on the time series trend for each section. The transverse cracking and joint faulting models were calibrated before calibrating the IRI model. Since all inputs to the IRI model could be obtained, it was calibrated outside Pavement-ME without rerunning it or using the CAT tool. IRI has a closed-form solution and does not require a standard error equation in Pavement-ME. The standard error for IRI is calculated using the standard error of its components.

4.7 CALCULATION OF DESIGN RELIABILITY

Pavement-ME uses a reliability-based design, as explained in Chapter 2. Reliability is added to the mean prediction to incorporate input or performance data variability. It is expressed as a

function of the predicted performance and derived using the predicted and measured performance data. A step-by-step approach to estimating the reliability of transverse cracking for rigid pavements is shown below as an example. A similar approach was used for the reliability of all other models except IRI in the Pavement-ME.

Step 1: All predicted and measured data points are grouped by creating bins on the predicted cracking. The number of data points in each group should be equivalent to reduce bias in the results.

Step 2: The average and standard deviation of measured and predicted cracking are computed for each group. The grouping is performed after finalizing the calibration coefficients (global or local) to obtain the predicted performance. Table 4-4 shows the number of data points, bin ranges, and descriptive statistics.

Table 4-4 Reliability analysis for transverse cracking in rigid pavements (example)

Cracking range (%)	No. of data points	Average Measured Cracking	Average Predicted Cracking	Standard dev. of Measured Cracking	Standard dev. of Predicted Cracking
0-0.5	46	0.84	0.54	0.86	0.29
0.5-2	31	1.41	1.35	1.51	0.25
2-5	44	3.53	3.13	3.76	0.72
5-10	29	1.45	12.18	8.93	1.58
10-50	12	15.06	26.52	14.96	1.22

Step 3: A relationship is determined between the standard deviation of the measured cracking on the y-axis and the average predicted cracking on the x-axis. Figure 4-4 shows the fit model to the grouped data in steps 1 and 2. Equation (4-25) shows the relationship between the standard deviation of the measured cracking and the average predicted cracking (when using the no-sampling technique).

$$s_{e(CRK)} = 1.3627(CRK)^{0.7473} \quad (4-25)$$

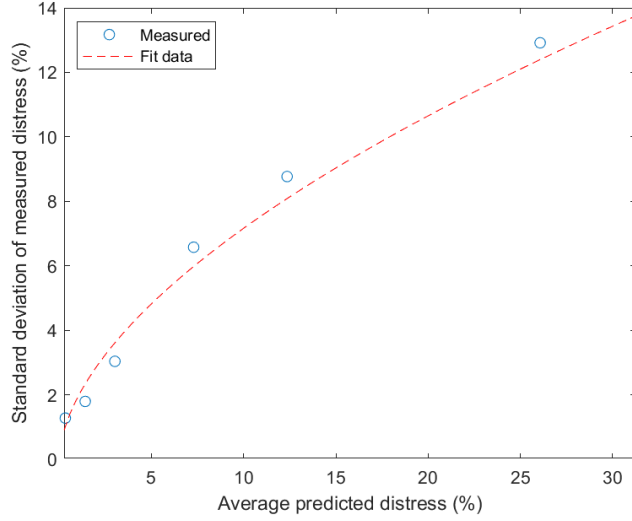


Figure 4-4 Fitting curve for the reliability of transverse cracking in rigid pavements (example)

Step 4: The reliability is calculated under the assumption that the error in prediction is approximately normally distributed towards the upper side of the mean distress. The predicted cracking can be adjusted to the desired reliability level using Equation (4-26)

$$C_r = C_{50} + S_e \times Z_{a/2} \quad (4-26)$$

where,

C_r = Predicted cracking at reliability r (%)

C_{50} = Predicted cracking at 50% reliability

S_e = Standard deviation of cracking, which can be estimated using Equation (4-25)

$Z_{a/2}$ = Standardized normal deviate (mean = 0; standard deviation = 1) at reliability r

Step 5: For the final step, the reasonableness of the model should be verified based on the actual measured data before using the reliability equation for design.

The reliability model for IRI is different from that of other models. Since it is a closed-form solution and the variances of different components of IRI are known, the reliability model for IRI is based on the variance analysis of its components. The basic assumption implies that the error in predicting IRI is roughly normally distributed. The total error includes input, repeatability, pure, and model errors. Overall, the IRI prediction error can be estimated by Equations (4-27) and (4-28).

$$IRI_{pe} = IRI_{meas} - IRI_{pred} \quad (4-27)$$

$$Var(IRI_{pe}) = Var(IRI_{meas}) + Var(IRI_{pred}) - 2R \times \sqrt{Var(IRI_{meas}) \times Var(IRI_{pred})} \quad (4-28)$$

where,

$Var(IRI_{pe})$ = Variance in prediction error for IRI (estimated from calibration results)

$Var(IRI_{meas})$ = Variance in measured IRI (estimated from field measurement)

$Var(IRI_{pred})$ = Variance in predicted IRI

R = Correlation coefficient between predicted and actual IRI

The variance in predicted IRI is the sum of the variance in inputs (cracking, spalling, faulting, and initial IRI) and the variance in model + pure error, as shown in Equation (4-29).

$$Var(IRI_{pred}) = Var(IRI_{INPUTS}) + Var(model + pure error) \quad (4-29)$$

The variance in inputs for the IRI model is shown in Equation (4-30).

$$Var(IRI_{INPUTS}) = Var_{IRIi} + C1^2 \times Var_{CRK} + C2^2 \times Var_{Spall} + C3^2 \times Var_{Fault} \quad (4-30)$$

where,

$Var(IRI_{INPUTS})$ = Variance in IRI due to measurement errors for each distresses and initial IRI (estimated from field measurements)

Var_{IRIi} = Variance in initial IRI

Var_{CRK} = Variance in transverse cracking

Var_{Spall} = Variance in joint spalling

Var_{Fault} = Variance in joint faulting

$C1, C2, C3$ = IRI model coefficients

Using Equations (4-28) to (4-30), $Var(model + pure error)$ can be determined and used to predict the standard deviation in IRI at any predicted value. The global standard error equations for each model are summarized in Table 4-5.

Table 4-5 Global calibration reliability equations for each distress and smoothness model

Pavement Type	Pavement performance prediction model	Standard error equation
Flexible pavements	Fatigue cracking (bottom-up)	$S_{e(bottom-up)} = 1.13 + \frac{13}{1 + e^{7.57 - 15.5 \times \log(FC_{Bottom} + 0.0001)}}$
	Fatigue cracking (top-down)	$S_{e(top-down)} = 0.3657 \times FC_{top} + 3.6563$
	Rutting	$S_{e(HMA)} = 0.24(\Delta_{HMA})^{0.8026} + 0.001$
		$S_{e(Base)} = 0.1477(\Delta_{Base})^{0.6711} + 0.001$
		$S_{e(SG)} = 0.1235(\Delta_{SG})^{0.5012} + 0.001$
	Transverse cracking	$s_e = 0.14 \times TC + 168$
Rigid pavements	IRI	Estimated internally by the software
	Transverse cracking	$S_{e(CRK)} = 3.5522(CRK)^{0.3415} + 0.75$
	Faulting	$S_{e(Fault)} = 0.07162(Fault)^{0.368} + 0.00806$
	IRI	Initial IRI $S_e = 5.4$ Estimated internally by the software

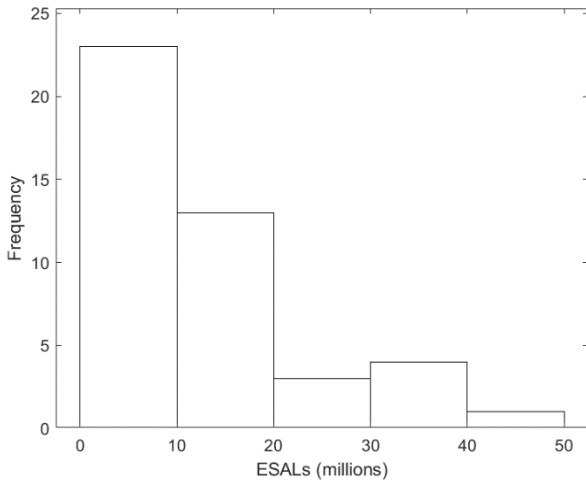
4.8 IMPACT OF CALIBRATION ON PAVEMENT DESIGN

Calibration aims to improve the Pavement-ME predictions and its usability for local conditions. The calibrated model will impact the local design practices. Additional flexible and rigid pavements (not part of the calibration) were designed to evaluate the impact of the locally calibrated models. The designs were based on calibrated model coefficients and standard error equations obtained using the least squares method. Forty-four (44) new flexible and 44 new rigid sections (JPCP) were designed in the Pavement-ME using the new calibrated models and the coefficients from the previous calibration effort (10). It is important to note that MDOT found the global coefficients more suitable than the local ones for actual designs. Therefore, the global coefficients were used for comparison in the case of rigid sections. Other design properties (base/subbase, subgrade, and climatic properties) were kept the same for flexible and rigid sections except for the traffic levels. These sections were also designed using the AASHTO93 design method. MDOT uses widened lane (lane width = 14 feet) sections for rigid pavements. The widened lane sections were designed as standard width (12 feet) by reducing the thicknesses by up to 1 inch from the final thickness. The lane width was kept at the standard width of 12 feet for flexible sections. Figure 4-5 shows the distribution of layer thicknesses (HMA and PCC), ESALs, and average annual MR for subgrade soil. The ESALs for flexible sections range from 1 to 41 million, whereas for rigid sections range from 1 to 64 million. The average annual MR for

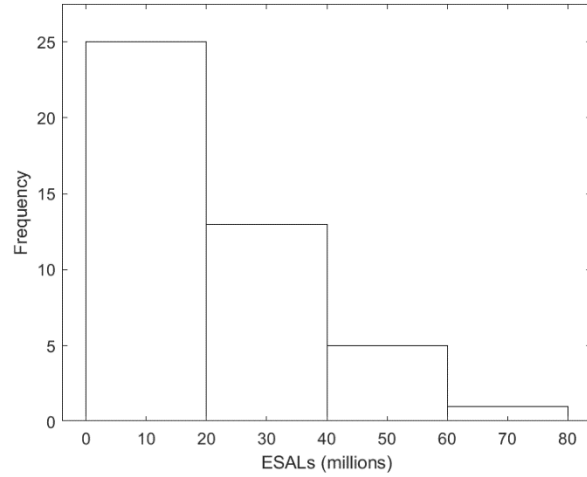
subgrade soil ranges from 3.7 to 6.5 ksi for flexible and rigid sections. Table 4-6 shows the number of sections in different categories.

All these flexible and rigid sections were designed in the Pavement-ME V2.6 at 95% design reliability and MDOT recommended thresholds. Table 4-7 shows the MDOT recommended threshold values for all distress types. Since the bottom-up cracking model was calibrated by combining the measured bottom-up and top-down cracking, the top-down cracking prediction was not used for design purposes. Moreover, MDOT does not have a formal design threshold for the new top-down cracking model.

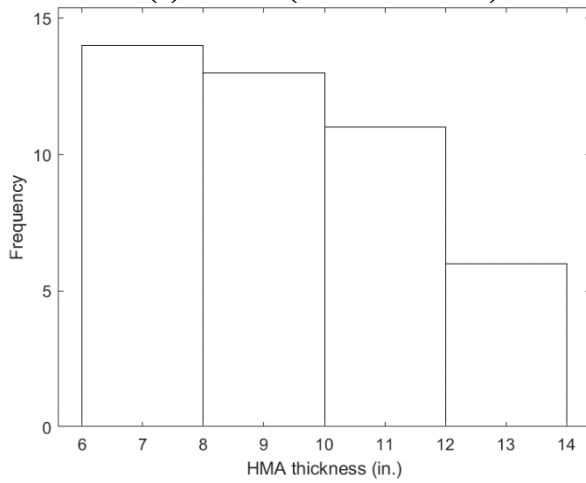
The design thicknesses were estimated to evaluate the differences between the newly calibrated model, previous calibrated model, and the AASHTO93 designs. Moreover, the critical design thicknesses were also identified separately for flexible and rigid pavements.



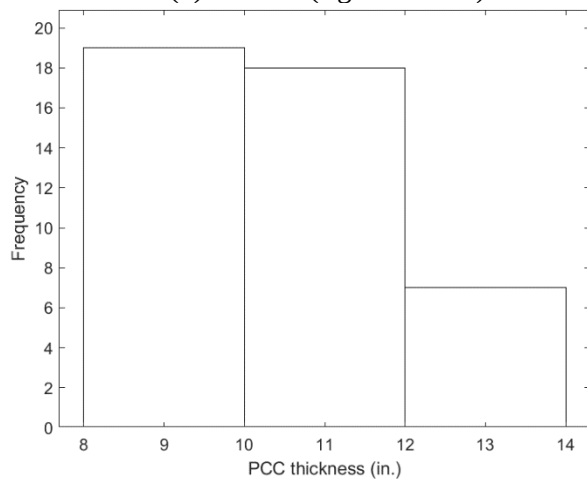
(a) ESALs (flexible sections)



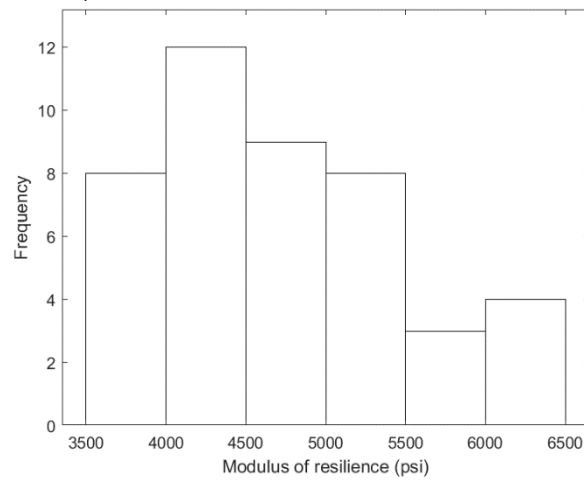
(b) ESALs (rigid sections)



(c) AASHTO93 HMA thickness (flexible sections)



(d) AASHTO93 PCC thickness (rigid sections)



(e) Subgrade MR (for both flexible and rigid sections)

Figure 4-5 Distribution of design inputs for the selected sections

Table 4-6 Selected sections for pavement design in different categories

Category	Description	# of sections
AASHTO soil classification	SP	6
	CL	18
	SC	8
	SM	3
	SC-SM	4
	SP-SM	5
MDOT region	Bay	4
	Grand	8
	Metro	8
	North	4
	Superior	3
	Southwest	8
	University	9
Classification	Freeway	28
	Non-freeway	16
Lane width (applicable to rigid sections only)	Widened (14 ft)	17
	Standard (12 ft)	27

Table 4-7 MDOT recommended design thresholds for Pavement-ME distress

Pavement type	Distress type	Threshold
Flexible pavements	Bottom-up cracking	20%
	Top-down cracking	NA
	Total rutting	0.5 inches
	Thermal cracking	1000 ft/mile
	IRI	172 in/mile
Rigid sections	Transverse cracking	15%
	Joint faulting	0.125 inches
	IRI	172 in/mile

4.9 SENSITIVITY ANALYSIS OF PAVEMENT-ME COEFFICIENTS

The sensitivity of the Pavement-ME transfer function coefficients is crucial in estimating the impact of each coefficient on the overall performance predictions. It is often not viable to calibrate all coefficients; therefore, only the sensitive ones can be estimated if the sensitivity of each coefficient is known. The sensitivity of the Pavement-ME transfer function coefficients was obtained using SSCs and NSI values for both flexible and rigid pavements. Moreover, they were compared to the NSI values from the literature (26). Four transfer functions were used for flexible pavements: bottom-up cracking, top-down cracking, total rutting, and IRI, whereas two transfer functions were used for rigid pavements: transverse cracking and IRI.

4.9.1 Sensitivity Using Normalized Sensitivity Index (NSI)

NSI has been typically used for this purpose and is defined as the percentage change of predicted distress relative to its global prediction caused by a given percentage change in the coefficient.

The NSI was calculated using Equation (4-31).

$$NSI = S_{ijk}^{DL} = \frac{\Delta Y_{ji}}{\Delta X_{ki}} \frac{X_{ki}}{Y_j} \quad (4-31)$$

where;

NSI = Normalized sensitivity index,

S_{ijk}^{DL} = Sensitivity index for input k , distress j , and at point i with respect to a given global prediction

ΔY_{ji} = Change in distress j around point i ($Y_{j,i+1} - Y_{j,i-1}$)

X_{ki} = Value of input X_k at point i

ΔX_{ki} = Change in input X_k around point i ($X_{k,i+1} - X_{k,i-1}$)

Y_j = Global prediction for distress j

The NSI values were also calculated to compare them with the results from SSCs. These calculations are based on the NCHRP 1-47 study (60) as shown in Equation (4-31). Ten sections, each from flexible and rigid pavements, were selected for NSI calculations. These sections exist in the MDOT PMS database, designed using the AASHTO93 design method. These sections are also part of the selected sections for calibration. It is essential to mention that for NSI calculation, each section was modeled in the Pavement-ME with the necessary design inputs (material, traffic, and climate). These inputs were obtained from construction records, job mix formulas, and other sources. Obtaining the design input is tedious and requires multiple data sources, unlike the calculation of SSCs, which does not require any data. The selected sections have a wide range of thicknesses and traffic. Tables 4-8 and 4-9 show the Pavement-ME inputs for flexible and rigid sections, respectively. Each section was initially run at the global values of transfer function coefficients at 50% reliability. Afterward, each coefficient (one at a time) was varied by -50%, -20%, 20%, and 50%, respectively, from the global values. The change in performance prediction was evaluated for differences in transfer function coefficients to calculate the NSI values.

Table 4-8 Design inputs for flexible sections used in NSI calculations

Section no.	HMA thickness (in.)	Base thickness (in.)	Subbase thickness (in.)	AADTT 2-way
1	8	6	18	2034
2	6.5	6	18	685
3	10.8	6	18	4315
4	4.3	6	12	201
5	5.5	6	15	859
6	5.5	6	18	959
7	14	16	8	6745
8	10.9	6	8	2065
9	8	4	18	354
10	6.5	6	18	313

Table 4-9 Design inputs for rigid sections used in NSI calculations

Section no.	PCC thickness (in.)	Base thickness (in.)	Subbase thickness (in.)	Dowel diameter (in.)	AADTT 2-way
1	11	4	14	1.5	7387
2	9.9	3.9	10	1.25	4825
3	12.2	3.9	10	1.5	12030
4	10.8	4	12	1.25	500
5	9.5	4	12	1.25	2758
6	10.8	6	12	1.5	10.8
7	12.5	16	0	1.5	12.5
8	11.7	4	10	1.5	11.7
9	11.3	3.9	12	1.5	11.3
10	11	4	10	1.5	11

While NSI can rank the coefficients based on their level of sensitivity, it does not provide information about any potential correlation between them or how accurately these can be estimated. Moreover, since NSI calculation requires distress data, its magnitude can change if the data source is changed; hence, the sensitivity ranking of the coefficients may vary (11).

4.9.2 Sensitivity Using Scaled Sensitivity Coefficient (SSC)

Unlike NSI calculation, SSCs do not require input data. SSCs were calculated for a continuous range of independent variables, and the results were visualized as SSC plots. The i_{th} sensitivity coefficient of a model, $\eta(x, \beta)$, where x is an independent variable, and β represents the parameter vector, is given by $X_i = \partial \eta / \partial \beta_i$ and indicates the magnitude of change of the response resulting from a small perturbation in the parameter β_i (64). An initial parameter value is required if the model is nonlinear in that parameter, i.e., $\partial \eta / \partial \beta_i = f(\beta_i)$, and requires an iterative solution using

any nonlinear regression algorithm (64). The parameter's SSC is the product of its sensitivity coefficient and the parameter itself, as shown in Equation (4-32).

$$X'_i = \beta_i \frac{\partial \eta}{\partial \beta_i} \quad (4-32)$$

where;

X'_i = Scaled sensitivity coefficient of the parameter i ,

β_i = Estimate of the i th parameter,

$\frac{\partial \eta}{\partial \beta_i}$ = i th sensitivity coefficient of the model w.r.t β_i .

Assume that a model $\eta(x, \beta)$ has two parameters, β_1 and β_2 . The sensitivity coefficients (X_i) and SSC (X'_i) for both parameters are estimated using the following equations [Equations (4-33) to (4-36)]. Suppose the parameters (β) have been estimated using any nonlinear regression algorithm, and the sensitivity coefficient matrix J is obtained. In that case, the SSC for either parameter can be approximated using Equations (4-37) and (4-38).

$$X_1 = \frac{\partial \eta}{\partial \beta_1} \approx \frac{\eta((1.001 * \beta_1), \beta_2) - \eta(\beta_1, \beta_2)}{0.001 * \beta_1} \quad (4-33)$$

$$X'_1 = \beta_1 \frac{\partial \eta}{\partial \beta_1} \approx \frac{\eta((1.001 * \beta_1), \beta_2) - \eta(\beta_1, \beta_2)}{0.001} \quad (4-34)$$

$$X_2 = \frac{\partial \eta}{\partial \beta_2} \approx \frac{\eta(\beta_1, (1.001 * \beta_2)) - \eta(\beta_1, \beta_2)}{0.001 * \beta_2} \quad (4-35)$$

$$X'_2 = \beta_2 \frac{\partial \eta}{\partial \beta_2} \approx \frac{\eta(\beta_1, (1.001 * \beta_2)) - \eta(\beta_1, \beta_2)}{0.001} \quad (4-36)$$

$$X'_1 \approx \beta_1 * J(:, 1) \quad (4-37)$$

$$X'_2 \approx \beta_2 * J(:, 2) \quad (4-38)$$

The SSC for a particular coefficient (say β_i) is calculated by differentiating the function w.r.t. β_i and multiplying it by β_i [as shown in Equation (4-32)]. Other coefficients except β_i are held constant. A similar approach is used to calculate SSCs for all other coefficients. The mathematical model (transfer function) can often be complicated, especially when differentiating the function. In that case, the SSCs can be approximated numerically to avoid errors in the analytical derivation. An example of the estimation of SSCs using the transverse cracking model [shown in Equation (4-39)] for rigid pavements.

$$CRK = \frac{1}{1 + C_4(DI_F)^{C_5}} \quad (4-39)$$

where,

CRK = Predicted fraction of bottom-up or top-down cracking

DI_F = Total fatigue damage (bottom-up or top-down)

C_4, C_5 = Transfer function coefficients

Denoting transverse cracking as a function of DI_F , C_4 , and C_5 [$CRK(DI_F, C_4, C_5)$], the sensitivity coefficient for C_4 (X_{C_4}) can be approximated as shown in Equation (4-40).

$$\frac{\partial CRK}{\partial C_4} = X_{C_4} \approx \frac{CRK(DI_F, C_4 + \delta, C_5) - CRK(DI_F, C_4, C_5)}{\delta \times C_4} \quad (4-40)$$

Here δ is a small quantity (a value of 0.001 was used). The SSC for C_4 (X'_{C_4}) can be approximated as shown in Equation (4-41).

$$\begin{aligned} C_4 \frac{\partial CRK}{\partial C_4} = X'_{C_4} &\approx \frac{CRK(DI_F, C_4 + \delta, C_5) - CRK(DI_F, C_4, C_5)}{\delta \times C_4} \\ &= \frac{CRK(DI_F, C_4 + \delta, C_5) - CRK(DI_F, C_4, C_5)}{\delta} \end{aligned} \quad (4-41)$$

The coefficient C_4 was changed by δ to get the first term of the numerator. The second term of the numerator is the transverse cracking at global values. Both these terms were evaluated at a continuous range of DI_F from 0 to 1. This provides a continuous set of X'_{C_4} for each value of DI_F . SSCs for C_5 (X_{C_5}) was calculated as shown in Equation (4-42). SSCs for each coefficient were plotted with DI_F in the same plot. A similar process was used for all other transfer functions.

$$\begin{aligned} C_5 \frac{\partial CRK}{\partial C_5} = X'_{C_5} &\approx \frac{CRK(DI_F, C_4, C_5 + \delta) - CRK(DI_F, C_4, C_5)}{\delta \times C_5} \\ &= \frac{CRK(DI_F, C_4, C_5 + \delta) - CRK(DI_F, C_4, C_5)}{\delta} \end{aligned} \quad (4-42)$$

The SSCs were calculated and plotted using MATLAB codes using one coefficient at a time and considering other coefficients as constant. A wide range of independent variables have been used since calculating SSCs is a forward problem without data.

4.10 CHAPTER SUMMARY

This chapter detailed the calibration approach used for each Pavement-ME prediction model. Transfer functions have been calibrated based on whether they calculate the distresses directly or based on cumulative damage. It also discusses the different resampling techniques and optimization methods. No sampling, bootstrapping, traditional split sampling, and repeated split sampling techniques have been used for calibration. For calibration validation, traditional and repeated split sampling were used. The calibration methods include the least squares and MLE. The process used for the MLE methodology is also outlined in this chapter. Reliability analysis is detailed, illustrating steps for estimating reliability equations for distress prediction, considering the transverse cracking as an example. Additionally, this chapter discusses the approach to assess the impact of calibration on pavement design based on thicknesses and critical distresses. Sensitivity analysis was conducted using Normalized Sensitivity Index (NSI) and Scaled Sensitivity Coefficients (SSCs), providing insights into the impact of model coefficients on performance predictions. These analyses facilitate the identification of sensitive coefficients crucial for accurate predictions and design decisions.

CHAPTER 5 - RESULTS AND DISCUSSION

The calibration process adjusts the Pavement-ME model parameters to match observed data better to ensure that the model outputs are reliable and valuable for pavement design. The Pavement-ME models' calibration process can be challenging because of their complexity and the large number of parameters involved. However, technological advancements and data collection methods have made the calibration process more efficient and effective. This chapter documents the results for calibration of each model, pavement design, and sensitivity of the Pavement-ME coefficients. Table 5-1 summarizes the calibration method used for each Pavement-ME model.

Table 5-1 Summary of calibration method for each Pavement-ME model

Pavement type	Pavement-ME model	Calibration method	
		MLE	Least squares
Flexible pavement	Bottom-up cracking: Option a		✓
	Bottom-up cracking: Option b	✓	✓
	Top-down cracking		✓
	Rutting (Method 1)		✓
	Rutting (Method 2)	✓	✓
	Thermal cracking		✓
	IRI	✓	✓
Rigid pavement	Transverse cracking	✓	✓
	Joint faulting		✓
	IRI	✓	✓

5.1 LOCAL CALIBRATION RESULTS FOR FLEXIBLE PAVEMENTS

This section presents the results for the local calibration of the bottom-up cracking, total rutting, and IRI models. Bottom-up cracking was calibrated using synthetic and observed data. It is important to note that bottom-up cracking using Option a, rutting using Method 1, top-down and thermal cracking models using observed data were calibrated using the least squares method only, as shown in Table 5-1. The calibration results for these models are shown in Table 5-2. These results correspond to the bootstrap resampling technique. The details of these model calibrations are shown in the Appendix.

Table 5-2 Summary of flexible pavement models calibrated using only the least squares method

Pavement-ME model	Local coefficient		Global model		Local model	
			SEE	Bias	SEE	Bias
Bottom-up cracking (Option a)	$C_1 = 0.2320$ $C_2 = 0.6998$ (hac < 5 in) $C_2 = (0.867 + 0.2583 * hac) * 0.2204$ (5 in <= hac <= 12 in) $C_2 = 0.8742$ (hac > 12 in)		8.30	-4.91	8.73	0.00
Top-down cracking	$K_{L1} = 64271618$ $K_{L2} = 0.90$ $K_{L3} = 0.09$ $K_{L4} = 0.101$ $K_{L5} = 3.260$ $C_1 = 0.30$ $C_2 = 1.155$ $C_3 = 1$		6.37	-2.36	5.59	1.60
Rutting (Method 1)	HMA	$\beta_{1r} = 0.148$ $\beta_{2r} = 0.7$ $\beta_{3r} = 1.3$	0.256	0.201	0.080	-0.013
	Base	$\beta_{s1} = 0.301$	0.042	0.038	0.009	-0.001
	Subgrade	$\beta_{sg1} = 0.070$	0.118	0.109	0.006	-0.000
Thermal cracking	$K = 0.85$		1225	-812	851	20

5.1.1 Calibration Using Synthetic Data

As mentioned in Chapter 4, exponentially distributed synthetic data was generated for bottom-up cracking with and without variability. Figure 5-1 shows the generated data distribution and different fitted probability distributions. The normal distribution legend in Figure 5-1 corresponds to the least squares method, while other distributions are used for the MLE method. The distribution is skewed for Figures 5-1(a) and 5-1(b) so that more data points are less than 5%, showing that the data is not normally distributed.

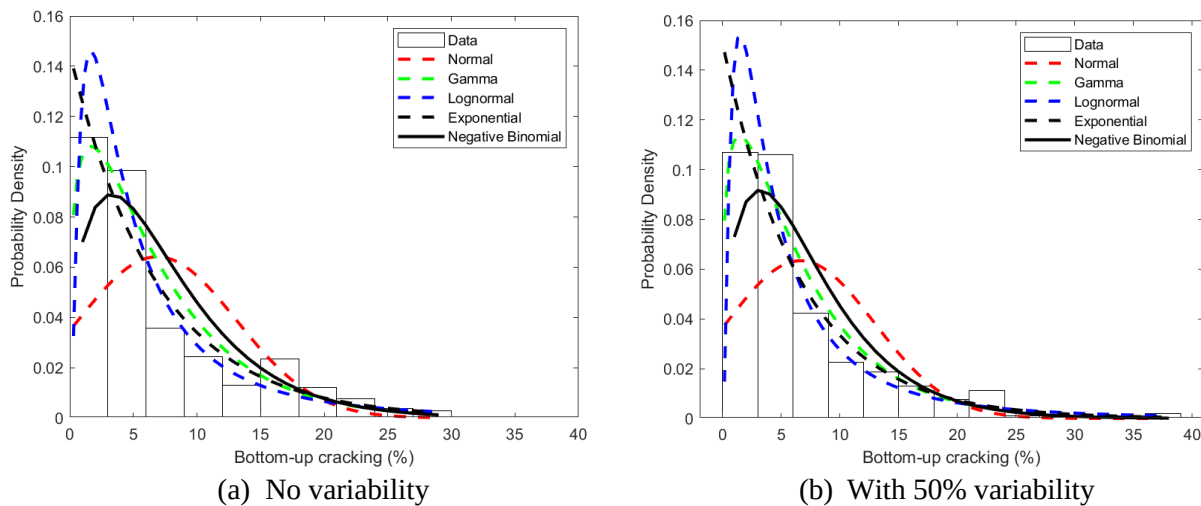


Figure 5-1 Distribution of synthetic data (bottom-up cracking in flexible pavements)

Both sets of generated data were calibrated using MLE and least squares methods, as well as the four mentioned sampling techniques. Table 5-3 summarizes both data sets' no-sampling and bootstrapping calibration results. As previously described, resampling techniques provide confidence intervals for the population parameter. Bootstrapping calibration results in Table 5-3 are the mean values. MLE provides better statistical parameters (SEE, bias, NLL, AIC, and BIC values) for all distributions, except for SEE values in the case of exponential distribution. The SEE values show the variability between predicted and measured data. The higher the SEE value, the more the dispersion along the line of equality. Pavement-ME uses a reliability-based design. Higher values of SEE imply higher reliability imposed over mean Pavement-ME predictions. Gamma distribution provides the best parameter estimates for the MLE method, followed by negative binomial distribution. It is worth mentioning that the parameter estimates (C_1 and C_2) using MLE are much closer to the assumed initial values than the estimates from the least squares method. Tables 5-4 and 5-5 summarize the results from split sampling and repeated split sampling techniques. Table 5-5 shows the mean values using the repeated split sampling method. Similar results are obtained, where MLE provides better statistical parameters than the least squares method. The gamma and negative binomial distribution for these validation results also offer optimum results with the SEE and bias values significantly lower than the least squares method. It is worth mentioning that the results from resampling techniques provide better parameter estimates, as can be seen in Tables 5-3 and 5-5.

Table 5-3 Summary of calibration results for synthetic data in flexible pavements

Calibration method	Distribution	With no variability					With 50% variability				
		SEE	Bias	NLL	AIC	BIC	SEE	Bias	NLL	AIC	BIC
No sampling	Normal	2.967	0.000	1190	2385	2393	7.305	0.000	1259	2522	2529
	Exponential	0.049	0.000	1040	2082	2086	5.690	0.000	1026	2055	2059
	Gamma	0.000	0.000	1032	2068	2076	2.584	0.000	1020	2045	2052
	Log normal	0.015	-0.007	1038	2079	2087	2.593	0.033	1028	2060	2068
	Negative binomial	0.002	-0.001	944	1891	1899	2.561	0.045	941	1886	1894
Bootstrapping	Normal	3.265	0.111	1235	2473	2481	4.282	0.143	1269	2542	2550
	Exponential	3.975	0.000	1015	2032	2036	4.986	0.000	1010	2022	2026
	Gamma	0.000	0.000	1008	2020	2028	2.553	0.000	1006	2016	2024
	Log normal	0.013	-0.007	1032	2068	2076	2.552	-0.148	1020	2044	2051
	Negative binomial	0.001	0.000	591	1186	1194	2.542	-0.001	683	1369	1377

Table 5-4 Summary of validation results using synthetic data in flexible pavements (Split sampling)

Data set	Distribution	With no variability					With 50% variability				
		SEE	Bias	NLL	AIC	BIC	SEE	Bias	NLL	AIC	BIC
Calibration set	Normal	0.231	0.000	798	1601	1608	6.962	0.000	870	1743	1750
	Exponential	1.707	0.000	718	1437	1441	2.838	0.000	710	1421	1425
	Gamma	0.001	-0.001	712	1427	1434	2.554	0.000	705	1413	1420
	Log normal	0.020	-0.010	712	1428	1435	2.569	0.037	707	1419	1426
	Negative binomial	0.042	0.018	652	1309	1316	2.515	0.040	652	1309	1316
Validation set	Normal	0.280	-0.041	352	708	713	8.453	1.057	390	784	790
	Exponential	2.028	-0.308	321	645	648	3.258	-0.175	316	634	637
	Gamma	0.001	-0.001	319	643	648	2.690	0.145	315	633	639
	Log normal	0.024	-0.013	319	643	648	2.703	0.196	315	633	639
	Negative binomial	0.050	0.025	291	586	591	2.689	0.129	289	581	587

Figure 5-2 compares both data sets' calibration results using MLE and least squares. Figures 5-2a and 5-2b show the propagation of bottom-up cracking with damage. The MLE predictions are closer to the synthetic measured data than the least squares predictions. This trend is more evident in Figure 5-2b, with 50% variability. Figures 5-2c, 5-2d, 5-2e, and 5-2f show the distribution of residuals (predicted – measured). Error distribution using MLE is less scattered and closer to zero. Moreover, it is closer to a normal distribution than the least squares method.

Table 5-5 Summary of validation results using synthetic data in flexible pavements (Repeated split sampling)

Data set	Distribution	With no variability					With 50% variability				
		SEE	Bias	NLL	AIC	BIC	SEE	Bias	NLL	AIC	BIC
Calibration set	Normal	4.436	0.203	917	1838	1845	5.146	0.253	958	1920	1927
	Exponential	3.825	0.000	722	1445	1449	4.851	0.000	720	1441	1445
	Gamma	0.000	0.000	719	1443	1450	2.522	0.000	719	1441	1448
	Log normal	0.023	-0.011	719	1442	1449	2.515	-0.059	716	1436	1443
	Negative binomial	0.001	0.000	473	950	957	2.508	0.006	501	1005	1012
Validation set	Normal	4.427	0.196	390	784	790	5.147	0.236	408	819	824
	Exponential	3.856	0.028	308	617	620	4.904	0.011	306	614	617
	Gamma	0.000	0.000	306	615	620	2.559	0.009	305	614	619
	Log normal	0.023	-0.011	306	616	621	2.544	-0.060	305	613	618
	Negative binomial	0.001	0.000	201	406	411	2.547	0.007	213	431	436

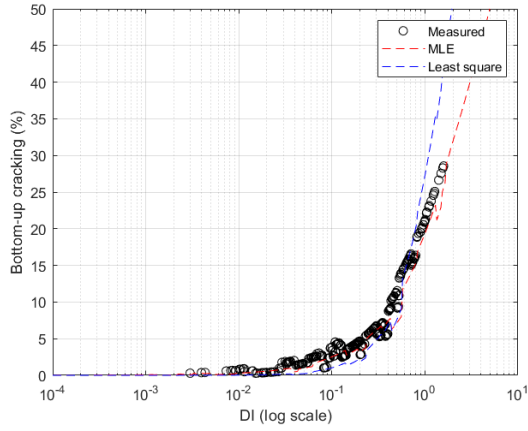
5.1.2 Calibration Using Observed Data

Based on the above process for synthetic data, the bottom-up cracking, total rutting, and IRI models were calibrated using MLE and least squares methods using observed data from field measurements. This observed data is obtained from MDOT's PMS database. Figure 5-3 shows

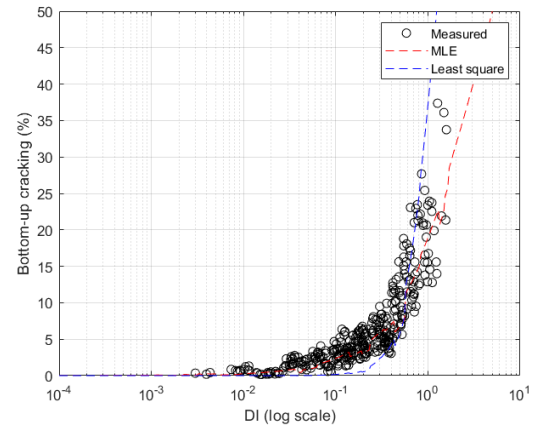
the distribution of observed data for different distresses and fitted distributions. Bottom-up cracking is the most skewed and non-normally distributed. Total rutting and IRI distributions are slightly skewed but closer to a normal distribution. As previously shown, resampling techniques provide better parameter estimates; therefore, bootstrapping and repeated split sampling results are presented.

Bottom-up cracking – Option b:

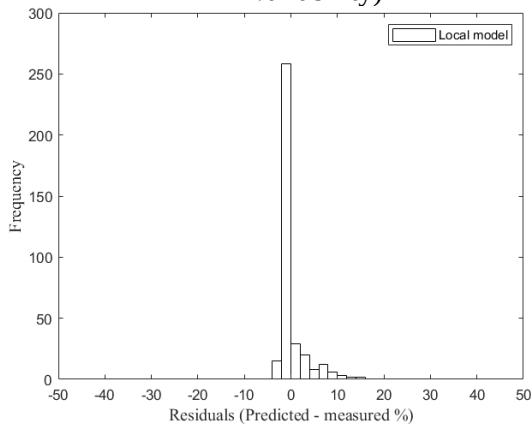
Table 5-6 summarizes bootstrapping and repeated split sampling results for bottom-up model calibration. MLE outperforms the least squares method with lower NLL, AIC, and BIC values for all distributions. The gamma distribution provides the best estimates for the MLE approach. Figure 5-4 shows the calibration results for bottom-up cracking using observed data using the bootstrapping technique for MLE (gamma distribution) and least squares methods. The predicted vs. measured plots show less MLE scatter than the least squares method.



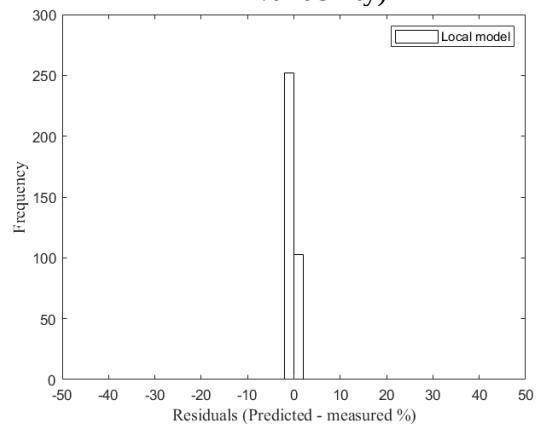
(a) Bottom-up cracking vs damage (with no variability)



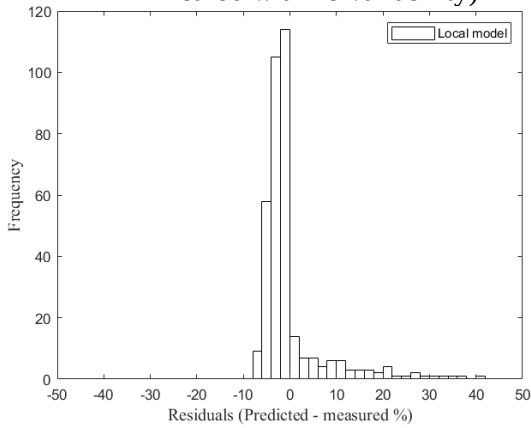
(b) Bottom-up cracking vs damage (with 50% variability)



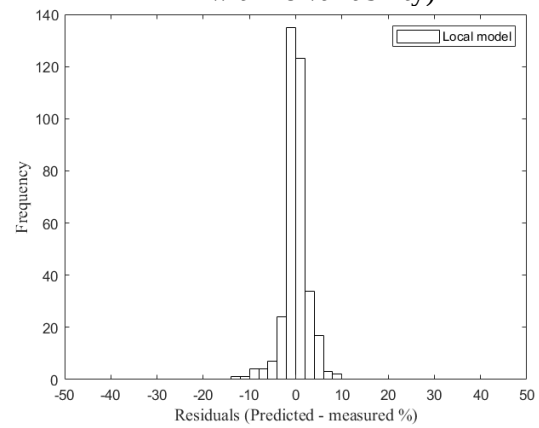
(c) Distribution of residuals (least squares method with no variability)



(d) Distribution of residuals (MLE method with no variability)



(e) Distribution of residuals (least squares method with 50% variability)



(f) Distribution of residuals (MLE method with 50% variability)

Figure 5-2 Calibration results for bottom-up cracking using synthetic data

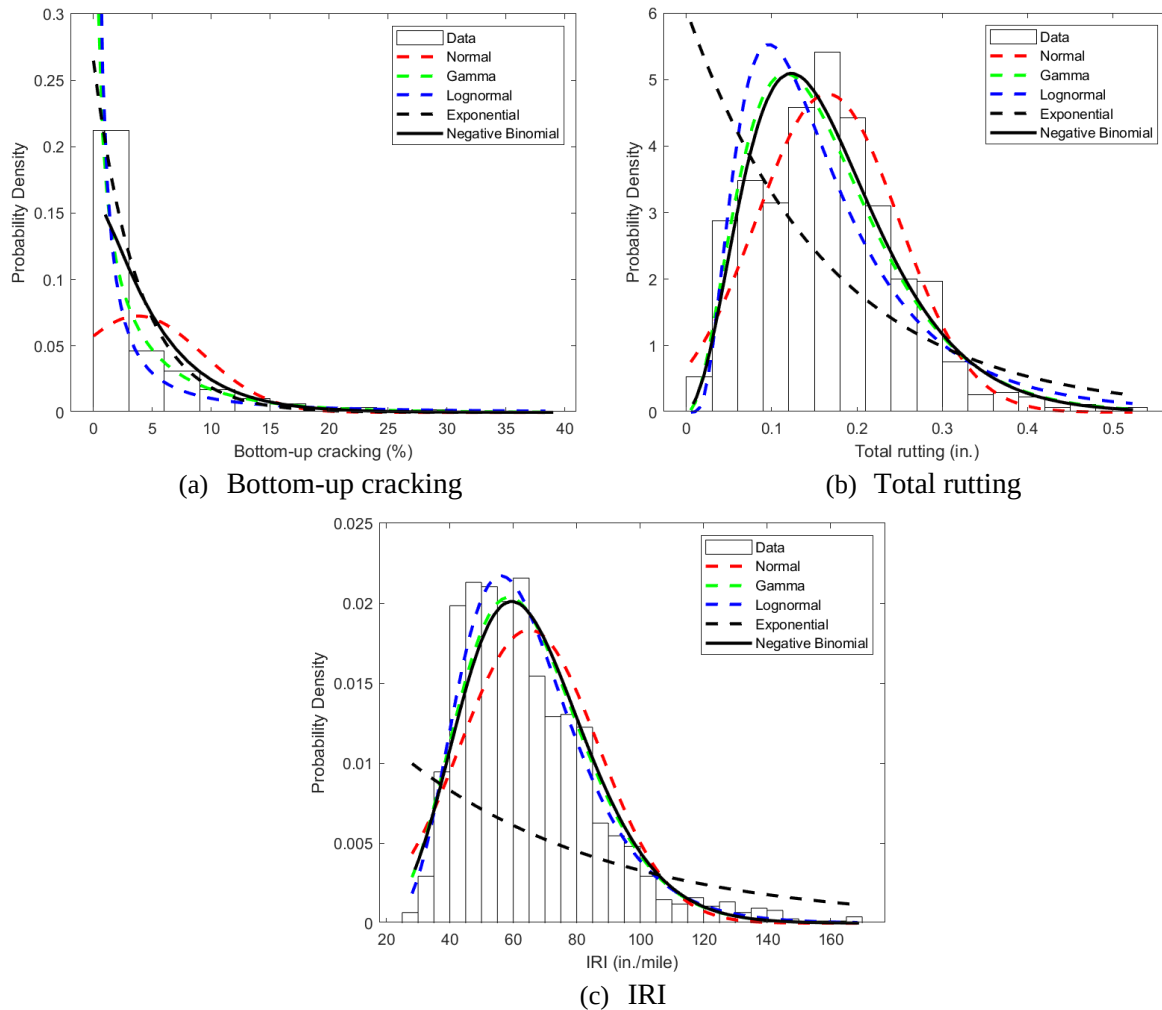


Figure 5-3 Distribution of observed data for flexible pavements

The distribution of residuals for MLE is also closer to zero. In Figures 5-4e and 5-4f, the red dashed line indicates the mean, the blue solid line shows the median, the red dashed line shows the 2.5th and 97.5th percentiles and the solid black line shows the cumulative distribution. Interestingly, the model parameters are normally distributed in the case of MLE, with the bias value consistently closer to zero.

Total rutting:

Table 5-7 shows the calibration results for the total rutting model. MLE shows better NLL, AIC, and BIC values for all MLE distributions compared to the least squares method. Gamma and negative binomial distributions provide the most feasible results using MLE. It also illustrates a bias-variance tradeoff where the SEE for gamma distribution is slightly higher than the least squares method but has a lower bias value. Figure 5-5 shows the calibration results using

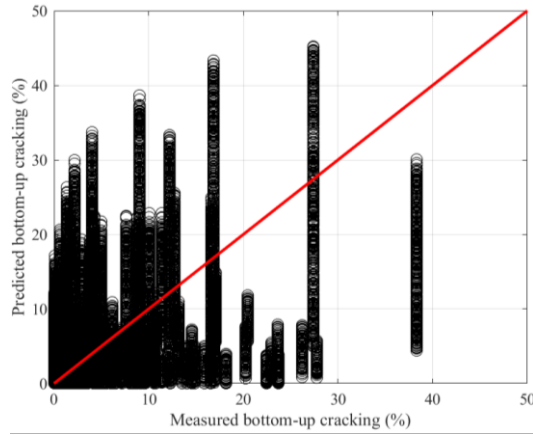
observed data for MLE (using gamma distribution) and the least squares methods. The predicted vs. measured plots show slightly less scatter for the MLE method. The residuals for MLE and least squares methods are comparable.

Table 5-6 Summary of calibration and validation results for observed data (Bottom-up cracking: Option b)

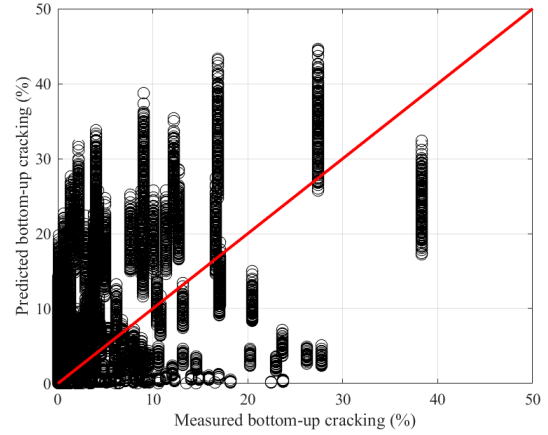
Calibration method	Distribution	SEE	Bias	C_1	C_2 (T<5 in.)	C_2 (T=5 to 12 in.)	NLL	AIC	BIC
Bootstrapping	Global	6.678	-3.769	1.310	2.159	1.000	3.4E+08	6.8E+08	6.8E+08
	Normal	6.114	0.052	0.221	0.716	0.234	1390	2784	2792
	Exponential	6.286	0.000	0.196	0.766	0.250	825	1652	1656
	Gamma	6.650	0.000	0.094	1.000	0.326	745	1495	1502
	Log normal	6.509	-0.160	0.112	0.974	0.318	759	1523	1530
	Negative binomial	5.517	0.424	0.467	0.133	0.043	870	1744	1752
Repeated split sampling (Calibration set)	Normal	6.183	0.021	0.210	0.745	0.243	975	1954	1961
	Exponential	6.239	0.000	0.206	0.733	0.239	579	1161	1164
	Gamma	6.692	0.000	0.095	0.997	0.326	525	1053	1060
	Log normal	6.503	-0.163	0.113	0.973	0.318	532	1068	1075
	Negative binomial	5.555	0.443	0.469	0.127	0.042	602	1208	1215
Repeated split sampling (Validation set)	Normal	6.224	0.043	0.210	0.745	0.243	420	844	849
	Exponential	6.281	0.025	0.206	0.733	0.239	248	497	500
	Gamma	6.792	0.064	0.095	0.997	0.326	248	500	505
	Log normal	6.597	-0.134	0.113	0.973	0.318	228	461	466
	Negative binomial	5.598	0.439	0.469	0.127	0.042	257	519	524

IRI:

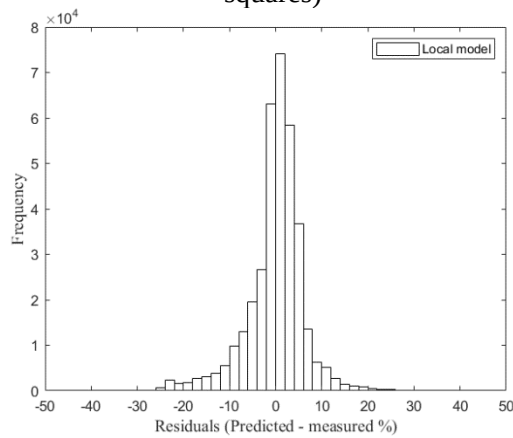
Table 5-8 shows the calibration results for the IRI model. The results from the MLE and least squares methods are comparable. The negative binomial distribution provides the best estimates among all distributions for the MLE method. Figure 5-6 shows the calibration results for MLE (negative binomial distribution) and the least squares methods. The predicted vs. measured plot shows slightly less scatter for MLE. The residual distribution between the MLE and least squares methods is comparable. In the case of IRI, the bias is consistently close to zero for the least squares method, showing that it is efficient for a robust calibration.



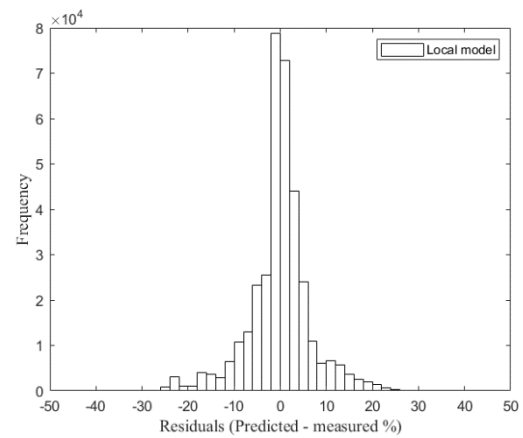
(a) Predicted vs. measured cracking (least squares)



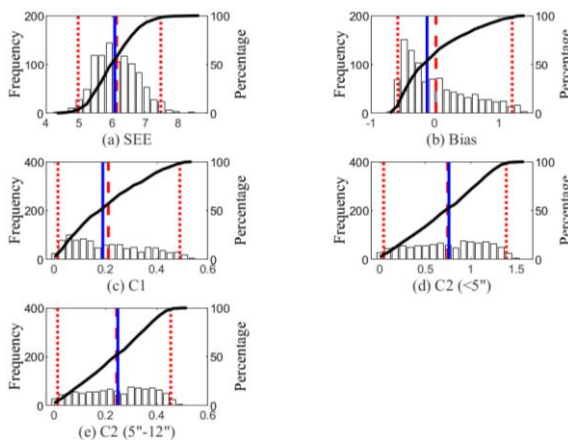
(b) Predicted vs. measured cracking (MLE)



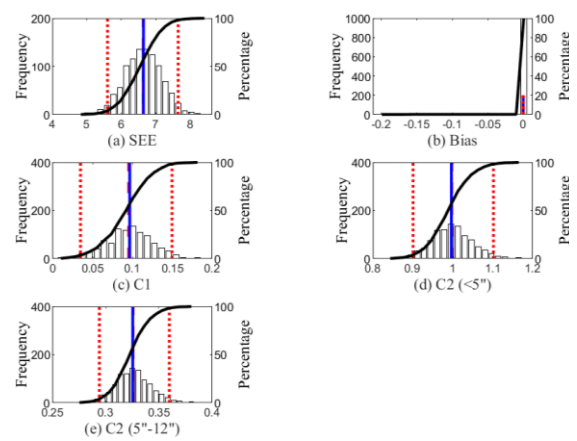
(c) Distribution of residuals (least squares)



(d) Distribution of residuals (MLE)



(e) Distribution of parameters (least squares)



(f) Distribution of parameters (MLE)

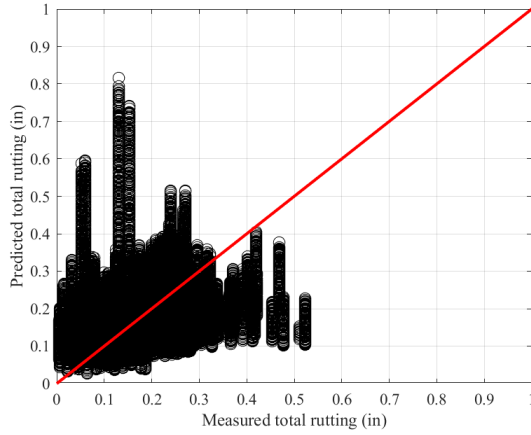
Figure 5-4 Calibration results for bottom-up cracking (Option b) using observed data

Table 5-7 Summary of calibration results for observed data (Total rutting)

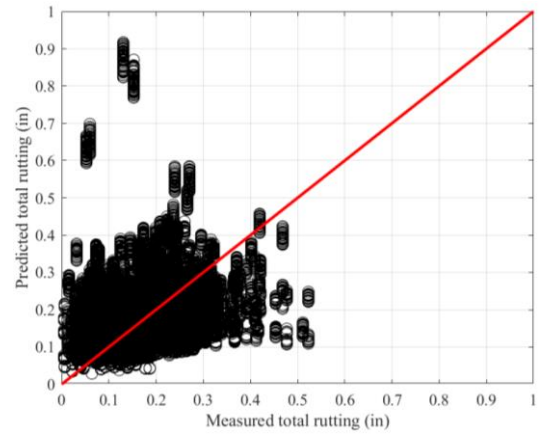
Calibration method	Distribution	SEE	Bias	β_{lr}	β_{sl}	β_{sq1}	NLL	AIC	BIC
Bootstrapping	Global	0.393	0.349	0.400	1.000	1.000	3784	7572	7582
	Normal	0.084	-0.008	0.144	0.839	0.523	2238	4481	4490
	Exponential	0.084	0.000	0.173	0.859	0.493	2146	4293	4298
	Gamma	0.096	0.000	0.129	0.163	0.396	2013	4031	4041
	Log normal	0.093	-0.012	0.102	0.158	0.490	2195	4394	4403
	Negative binomial	0.079	0.003	0.062	0.879	0.559	2230	4464	4474
Repeated split sampling (Calibration set)	Normal	0.078	0.000	0.028	1.185	0.634	1948	3901	3909
	Exponential	0.085	0.000	0.071	0.835	0.490	1503	3007	3012
	Gamma	0.096	0.000	0.124	0.160	0.432	1412	2827	2836
	Log normal	0.085	0.000	0.071	0.835	0.490	1503	3007	3012
	Negative binomial	0.079	0.003	0.063	0.869	0.558	1562	3127	3136
Repeated split sampling (Validation set)	Normal	0.080	-0.015	0.028	1.185	0.634	840	1683	1691
	Exponential	0.085	0.000	0.071	0.835	0.490	643	1289	1292
	Gamma	0.095	0.000	0.124	0.160	0.432	607	1219	1226
	Log normal	0.085	0.000	0.071	0.835	0.490	643	1289	1292
	Negative binomial	0.080	0.003	0.063	0.869	0.558	797	1597	1604

Table 5-8 Summary of calibration results for observed data (IRI - Flexible)

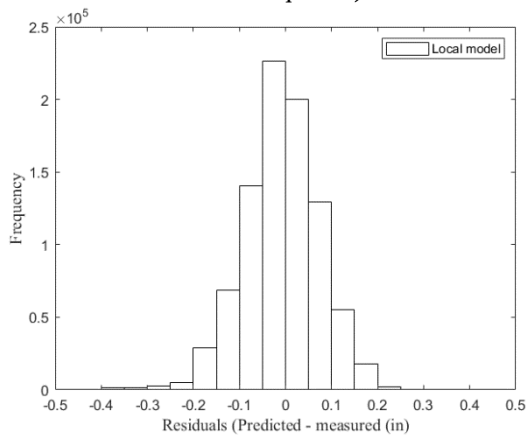
Calibration method	Distribution	SEE	Bias	C_1	C_2	C_3	C_4	NLL	AIC	BIC
Bootstrapping	Global	22.210	14.306	40.000	0.400	0.008	0.015	7368	14740	14751
	Normal	16.246	-0.630	41.486	0.433	0.006	0.0042	16996	33997	34007
	Exponential	16.406	0.008	43.033	0.485	0.007	0.0042	7773	15547	15553
	Gamma	18.943	1.273	40.022	0.312	0.020	0.0001	6631	13267	13277
	Log normal	18.573	0.593	40.026	0.195	0.019	0.00005	6590	13183	13194
	Negative binomial	15.606	-0.516	41.727	0.259	0.005	0.00617	7745	15493	15504
Repeated split sampling (Calibration set)	Normal	15.866	0.167	48.841	0.327	0.006	0.005	4948	9900	9910
	Exponential	16.419	0.000	43.485	0.516	0.006	0.0041	5444	10891	10896
	Gamma	18.951	1.280	40.038	0.324	0.019	0.000	4646	9297	9306
	Log normal	18.546	0.599	40.017	0.202	0.019	0.00002	4615	9235	9245
	Negative binomial	15.671	-0.512	41.714	0.261	0.005	0.00615	5425	10853	10863
Repeated split sampling (Validation set)	Normal	15.935	0.157	48.841	0.327	0.006	0.0051	2118	4241	4249
	Exponential	16.433	-0.006	43.485	0.516	0.006	0.0041	2329	4660	4664
	Gamma	19.034	1.288	40.038	0.324	0.019	0.000	1988	3980	3989
	Log normal	18.623	0.586	40.017	0.202	0.019	0.00002	1975	3953	3962
	Negative binomial	15.684	-0.502	41.714	0.261	0.005	0.00615	2009	4023	4031



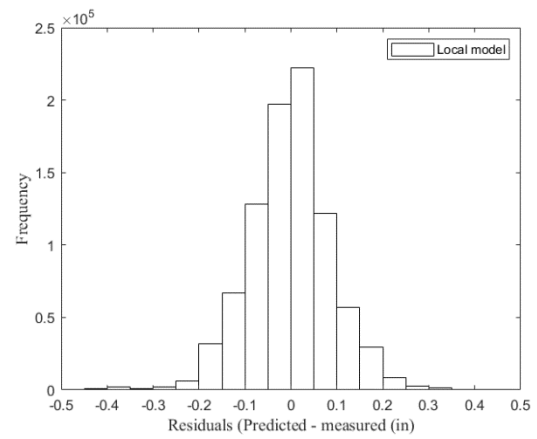
(a) Predicted vs. measured rutting (least squares)



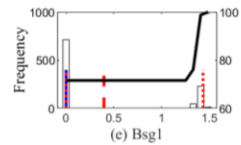
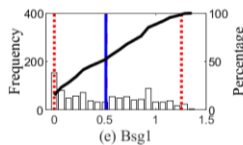
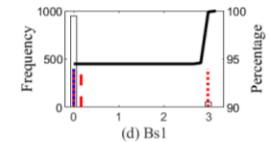
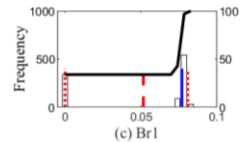
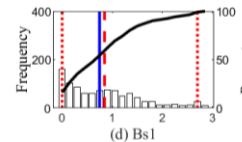
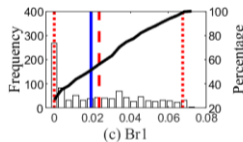
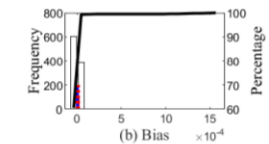
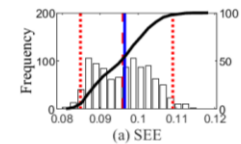
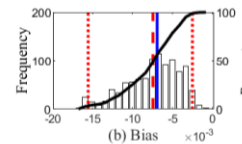
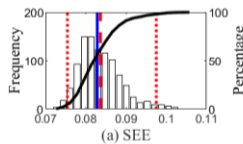
(b) Predicted vs. measured rutting (MLE)



(c) Distribution of residuals (least squares)



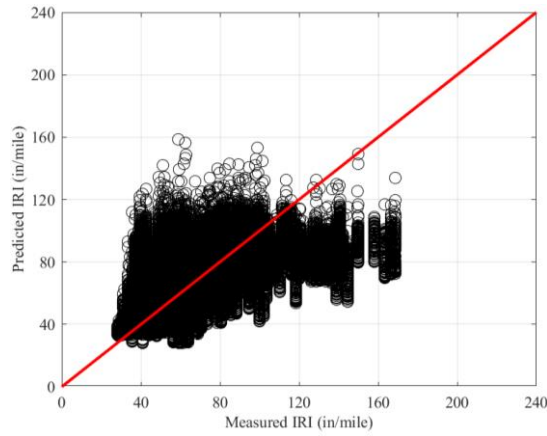
(d) Distribution of residuals (MLE)



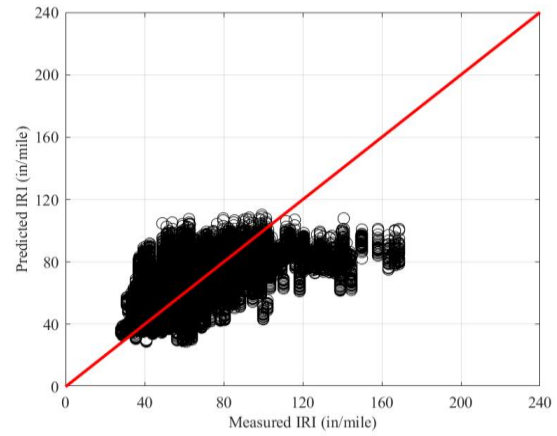
(e) Distribution of parameters (least squares)

(f) Distribution of parameters (MLE)

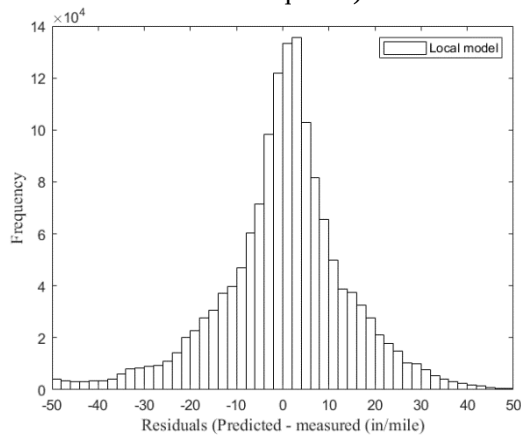
Figure 5-5 Calibration results for total rutting using observed data



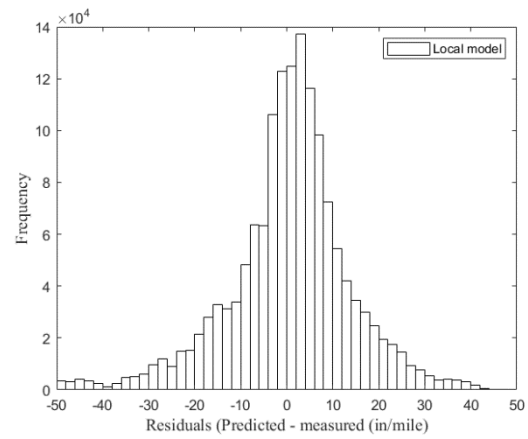
(a) Predicted vs. measured IRI (least squares)



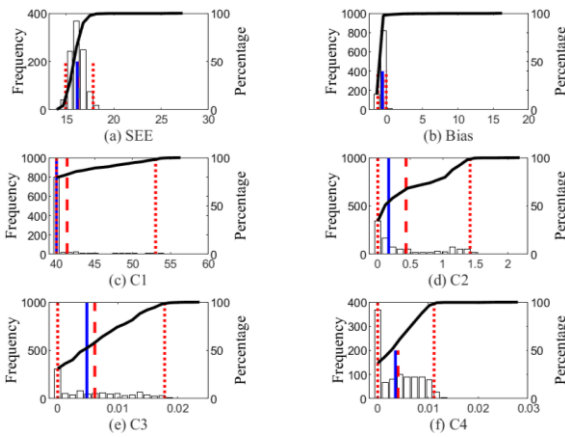
(b) Predicted vs. measured IRI (MLE)



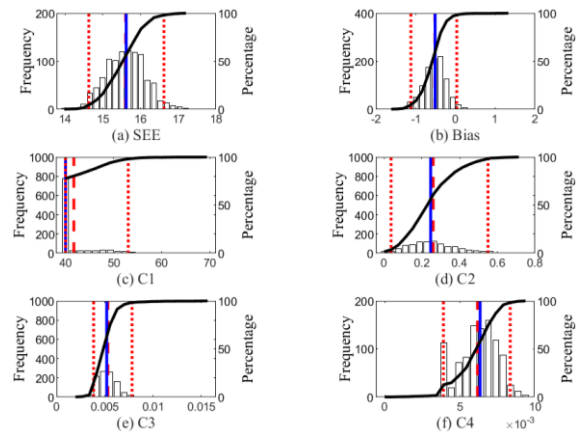
(c) Distribution of residuals (least squares)



(d) Distribution of residuals (MLE)



(e) Distribution of parameters (least squares)



(f) Distribution of parameters (MLE)

Figure 5-6 Calibration results for flexible IRI using observed data

5.2 LOCAL CALIBRATION RESULTS FOR RIGID PAVEMENTS

This section presents the results for the local calibration of the transverse cracking and IRI models. It is important to note that the joint faulting model was calibrated using the least squares method only, as shown in Table 5-1. Table 5-9 shows the calibration results for the joint faulting model, the details of which are shown in the Appendix.

Table 5-9 Summary of rigid pavement models calibrated using only the least squares method

Pavement-ME model	Local coefficient	Global model		Local model	
		SEE	Bias	SEE	Bias
Joint faulting	$C_1 = 0.8$ $C_2 = 1.3889$ $C_3 = 0.00217$ $C_4 = 0.00444$ $C_5 = 250$ $C_6 = 0.2$ $C_7 = 7.3$ $C_8 = 400$	0.06	0.01	0.03	0.00

5.2.1 Calibration Using Synthetic Data

Transverse cracking data was exponentially generated to study the effectiveness of using MLE with different conditions and distributions. Figure 5-7 shows the generated data with different fitted distributions. The normal distribution legend in Figure 5-7 corresponds to the least squares method.

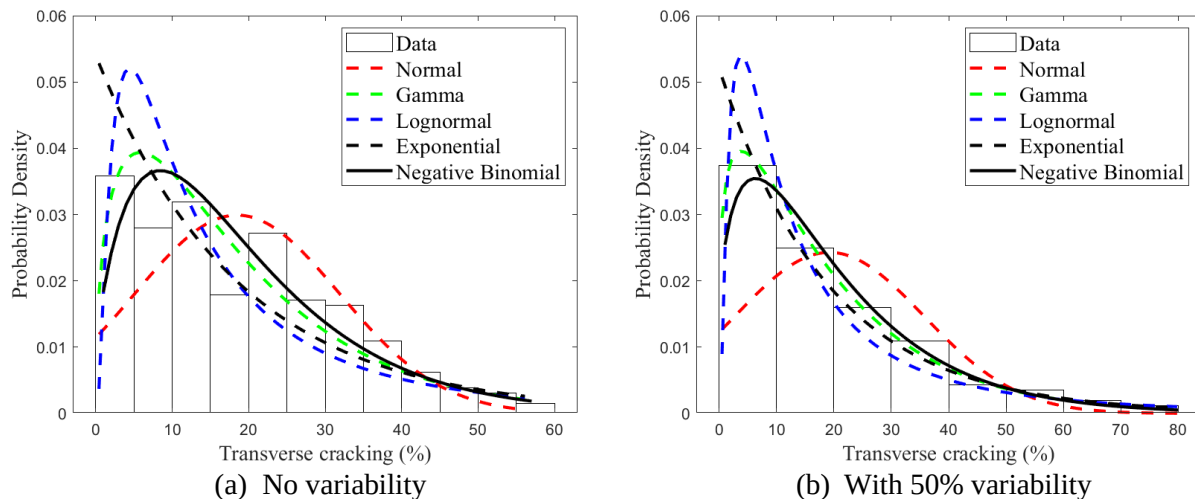


Figure 5-7 Distribution of synthetic data (transverse cracking in rigid pavements)

Calibration was performed using the least squares method and MLE for all four mentioned calibration and validation approaches: no sampling, bootstrapping, split sampling, and repeated split sampling. Table 5-10 summarizes calibration results using no sampling and bootstrapping methods. The mean values for the parameters are shown for the bootstrapping results in Table 5-10. The least squares method is denoted as Normal in Table 5-10. All distributions using MLE perform better than the least squares method in terms of NLL, AIC, and BIC. Except for exponential distribution, other distributions perform better regarding SEE values. Exponential and gamma distributions have lower bias than the least squares method. The gamma distribution is the most suitable distribution for this synthetic data. Compared to the least squares method, it provides better results for all parameters (SEE, bias, NLL, AIC, and BIC values).

A similar trend can be observed in the validation results. Tables 5-11 and 5-12 summarize the validation results using split sampling and repeated split sampling, respectively. The mean values for the parameters are shown for repeated split sampling results in Table 5-12. Gamma distribution provides better results than least squares regarding SEE, bias, NLL, AIC, and BIC values. This is more evident in resampling approaches. It is also a helpful illustration of the bias-variance tradeoff.

Table 5-10 Summary of calibration results for synthetic data in rigid pavements

Calibration method	Distribution	With no variability					With 50% variability				
		SEE	Bias	NLL	AIC	BIC	SEE	Bias	NLL	AIC	BIC
No sampling	Normal	3.586	-0.492	1040	2084	2092	6.996	-0.099	1081	2166	2173
	Exponential	0.705	0.000	1007	2016	2019	7.267	0.000	995	1992	1996
	Gamma	0.001	0.000	980	1965	1972	6.461	0.000	989	1982	1989
	Log normal	0.040	-0.023	1021	2046	2054	6.467	-0.369	1024	2053	2060
	Negative binomial	0.001	0.001	922	1847	1855	6.491	-0.051	920	1844	1851
Bootstrapping	Normal	2.315	-0.234	1016	2036	2043	7.159	-0.407	1075	2154	2161
	Exponential	4.628	0.000	993	1988	1992	8.515	0.000	995	1991	1995
	Gamma	0.001	0.000	979	1962	1969	6.422	0.000	988	1980	1987
	Log normal	0.036	-0.020	1019	2042	2049	6.454	-0.377	1023	2049	2056
	Negative binomial	0.017	-0.004	563	1131	1138	6.465	-0.107	725	1454	1462

Table 5-11 Summary of validation results using synthetic data in rigid pavements (Split sampling)

Data set	Distribution	With no variability					With 50% variability				
		SEE	Bias	NLL	AIC	BIC	SEE	Bias	NLL	AIC	BIC
Calibration set	Normal	0.000	0.000	703	1409	1416	6.314	-0.103	726	1455	1461
	Exponential	1.328	0.000	690	1381	1384	6.435	0.000	687	1376	1379
	Gamma	0.001	0.000	682	1367	1374	6.334	0.000	682	1368	1375
	Log normal	0.051	-0.028	709	1422	1428	6.317	-0.250	706	1417	1423
	Negative binomial	0.000	0.000	638	1280	1286	6.319	0.067	636	1276	1282
Validation set	Normal	0.000	0.000	303	610	614	7.065	-1.324	330	663	668
	Exponential	1.429	-0.188	304	609	611	7.690	-1.394	308	617	619
	Gamma	0.001	-0.001	298	600	605	6.855	-1.151	307	617	622
	Log normal	0.057	-0.034	312	627	632	7.037	-1.460	318	640	644
	Negative binomial	0.000	0.000	284	572	576	7.125	-1.174	284	572	577

Table 5-12 Summary of validation results using synthetic data in rigid pavements (Repeated split sampling)

Data set	Distribution	With no variability					With 50% variability				
		SEE	Bias	NLL	AIC	BIC	SEE	Bias	NLL	AIC	BIC
Calibration set	Normal	2.809	-0.287	726	1457	1463	7.139	0.005	743	1490	1496
	Exponential	4.996	0.000	697	1397	1400	7.989	0.000	702	1407	1410
	Gamma	0.000	0.000	711	1426	1432	6.199	0.000	694	1391	1398
	Log normal	0.052	-0.029	711	1427	1433	5.891	-0.535	722	1448	1454
	Negative binomial	0.034	-0.016	465	934	940	6.410	0.035	512	1028	1035
Validation set	Normal	2.828	-0.288	310	624	629	7.187	-0.036	318	640	645
	Exponential	5.052	-0.006	298	598	600	8.071	-0.017	301	603	606
	Gamma	0.000	0.000	304	611	616	6.265	-0.016	296	597	601
	Log normal	0.052	-0.030	304	611	616	6.020	-0.522	308	620	625
	Negative binomial	0.035	-0.016	199	402	407	6.528	0.016	219	443	448

The gamma distribution is most suitable for MLE and performs better than least squares estimates. Figure 5-8 shows the calibration results using the least squares method and MLE using a gamma distribution. The MLE predictions are closer to the measured data points (synthetic data), whereas the distribution of residuals shows a low scatter. The mean residual value is the model bias, whereas the spread of residuals represents the SEE. The mean SEE and bias values for the gamma distribution are 0.001 and 0.000, whereas, for the least squares method, they are 2.315 and -0.234, respectively, using bootstrap resampling on synthetic data with no variability. The mean bias value for the gamma distribution remains 0.000, whereas for the least squares, it

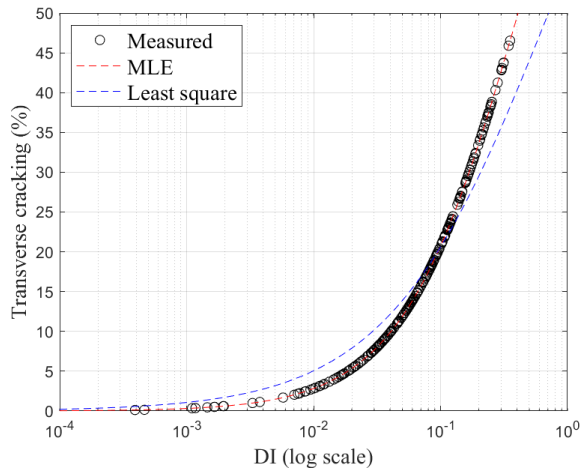
is -0.407, using the bootstrap resampling on synthetic data with 50% variability. This shows the robustness of the MLE method for data with different variabilities.

5.2.2 Calibration Using Observed Data

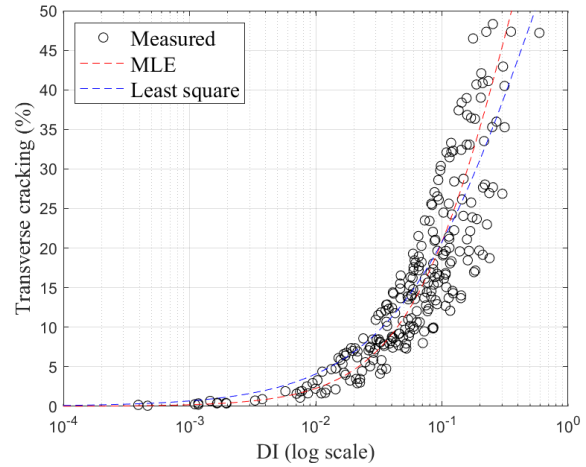
MLE and least squares methods were used to calibrate transverse cracking and IRI transfer functions using observed data obtained from MDOT's PMS data. Sections used for transverse cracking and IRI may differ, as the measured performance trends differ for both. Figure 5-9 shows the observed data distribution with different fitted distributions. Figure 5-9 shows that the transverse cracking data is skewed and non-normally distributed. IRI, on the other hand, is closer to a normal distribution.

Transverse cracking:

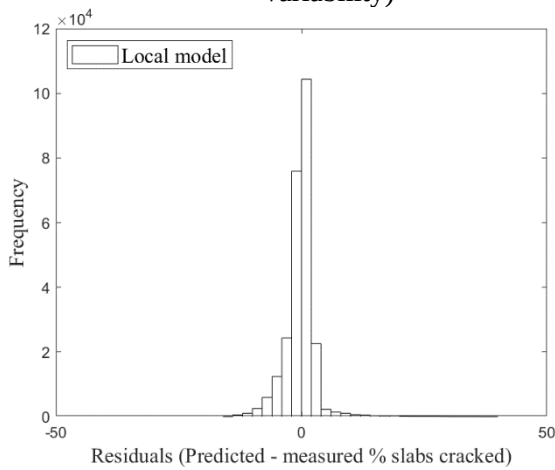
Table 5-13 summarizes the calibration and validation results for transverse cracking. Results for only the resampling approaches (bootstrapping and repeated split sampling) have been shown for brevity. The mean values for the parameters are shown in Table 5-13. MLE using gamma distribution provides the most feasible results with lower parameters (SEE, bias, NLL, AIC, and BIC) than the least squares. A similar trend is observed in the validation results using repeated split sampling (Table 5-13), where the MLE results show better validation parameters than the least squares method. Figure 5-10 shows the calibration results (for bootstrapping) for the least squares method and MLE using a gamma distribution. The predicted vs. measured transverse cracking shows a lower scatter for MLE. The mean bias value for the least squares is -0.410, whereas, for MLE, it is 0.000, using bootstrap results. The SEE values between the least squares and MLE are comparable. Also, the bias distribution for MLE is close to zero, illustrating the robustness of the MLE method. The lower and upper 95th percent confidence limits for the least squares are -0.932 and -0.025, whereas for the MLE, they are -0.001 to 0.001. This shows that MLE consistently has no bias for all 1000 bootstrap samples. Figure 5-10 (e) and (f) show the distribution of each bootstrap sample's SEE, bias, and transfer function coefficients. Bootstrap is used for 1000 resamples with replacement. A different set of parameters are obtained for each sample. These plots provide a distribution of parameters, and the mean value can be used as a more reliable estimate.



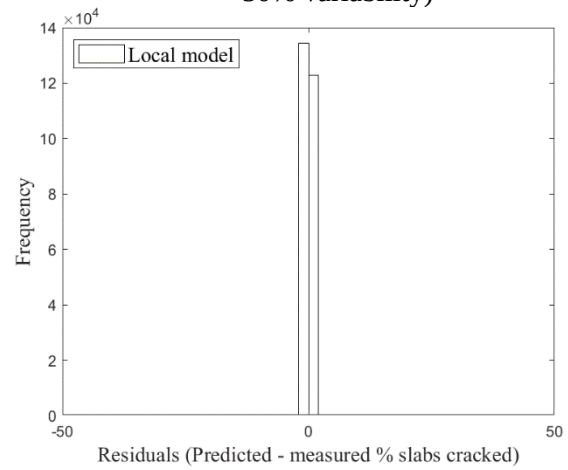
(a) Transverse cracking vs damage (with no variability)



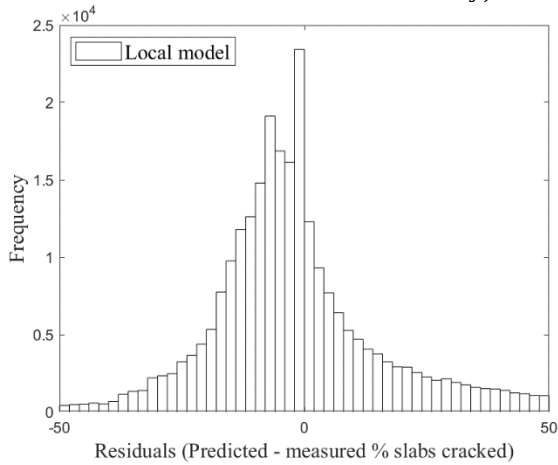
(b) Transverse cracking vs damage (with 50% variability)



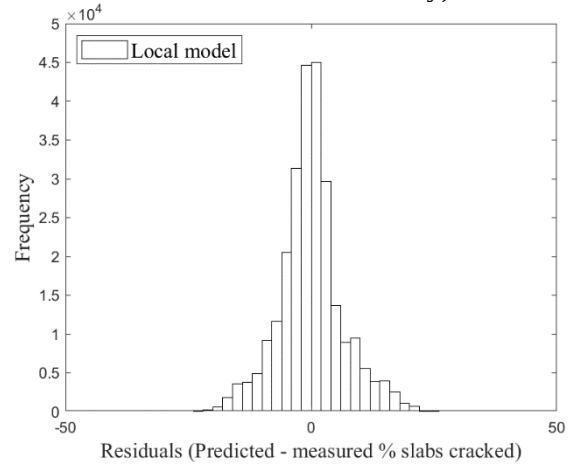
(c) Distribution of residuals (least squares method with no variability)



(d) Distribution of residuals (MLE method with no variability)



(e) Distribution of residuals (least squares method with 50% variability)



(f) Distribution of residuals (MLE method with 50% variability)

Figure 5-8 Calibration results for transverse cracking using synthetic data

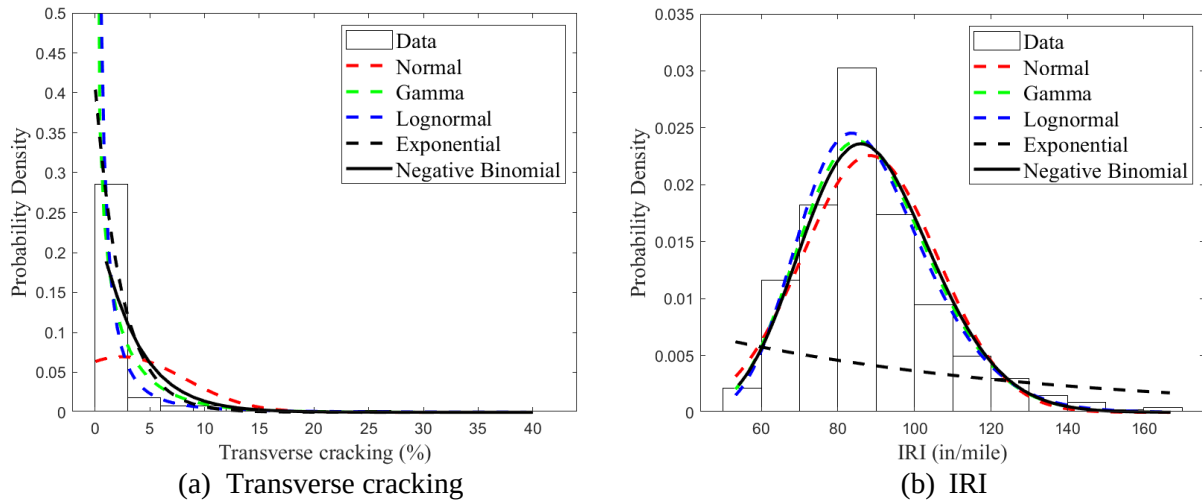


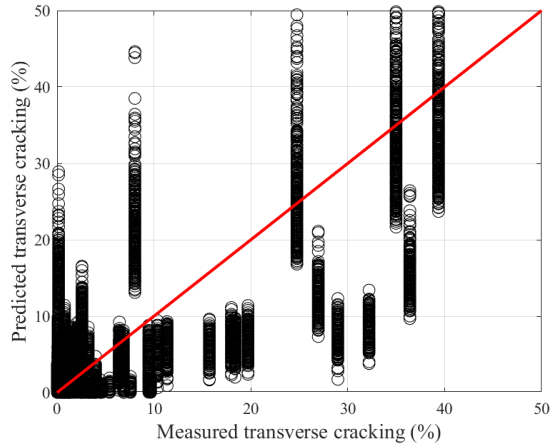
Figure 5-9 Distribution of observed data for rigid pavements

Table 5-13 Summary of calibration and validation results for observed data (Transverse cracking)

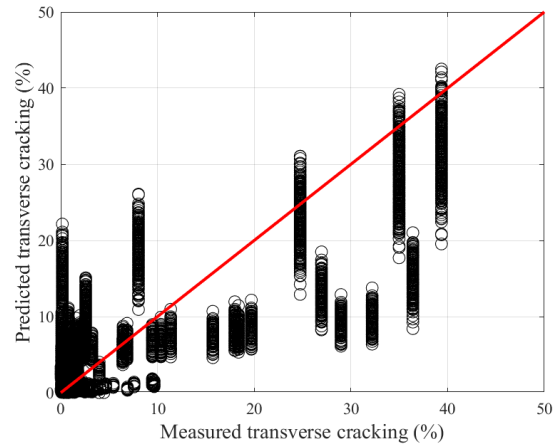
Calibration method	Distribution	SEE	Bias	C_4	C_5	NLL	AIC	BIC
Bootstrapping	Global	5.994	-2.390	0.52	-2.17	26051	52105	52112
	Normal	4.022	-0.410	0.476	-0.962	854	1713	1720
	Exponential	4.218	0.000	1.071	-0.708	484	970	973
	Gamma	3.984	0.000	0.668	-0.76	439	882	890
	Log normal	4.363	-0.578	1.406	-0.654	389	783	790
	Negative binomial	4.812	-0.166	4.563	-0.369	467	938	945
Repeated split sampling (Calibration set)	Normal	4.074	-0.411	0.467	-0.963	598	1200	1207
	Exponential	4.225	0.000	1.091	-0.682	340	682	686
	Gamma	4.038	0.000	0.650	-0.761	309	622	628
	Log normal	4.359	-0.577	1.406	-0.652	272	548	555
	Negative binomial	4.883	-0.184	4.704	-0.363	327	658	665
Repeated split sampling (Validation set)	Normal	4.129	-0.404	0.467	-0.963	270	543	548
	Exponential	4.252	0.018	1.091	-0.682	146	294	297
	Gamma	4.124	0.023	0.650	-0.761	132	268	273
	Log normal	4.374	-0.567	1.406	-0.652	118	240	245
	Negative binomial	4.900	-0.223	4.704	-0.363	142	287	292

IRI:

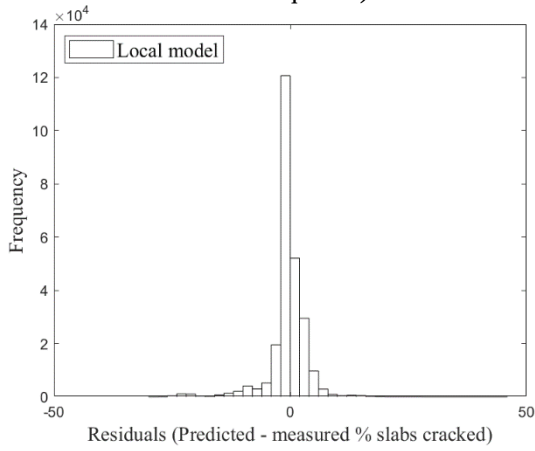
Table 5-14 summarizes the calibration results for IRI using the least squares and MLE methods. Table 5-14 shows the mean values for the parameters obtained using bootstrap resampling. MLE using negative binomial shows the most feasible results among all distributions. Interestingly, the least squares method shows satisfactory calibration and validation results, especially with lower SEE and bias values than MLE results.



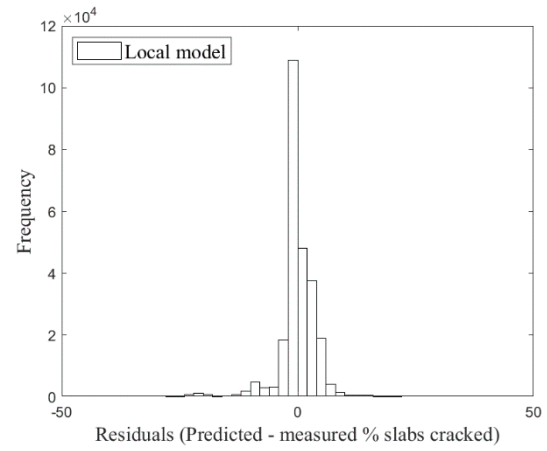
(a) Predicted vs. measured cracking (least squares)



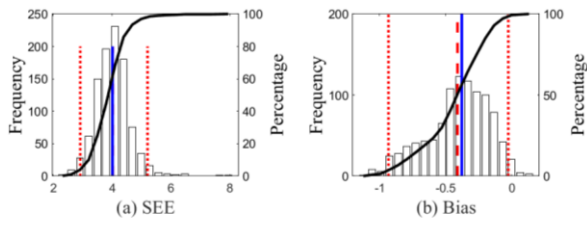
(b) Predicted vs. measured cracking (MLE)



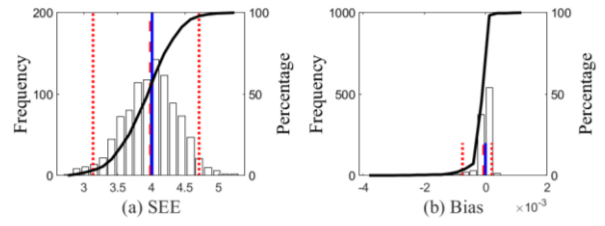
(c) Distribution of residuals (least squares)



(d) Distribution of residuals (MLE)



(e) Distribution of parameters (least squares)



(f) Distribution of parameters (MLE)

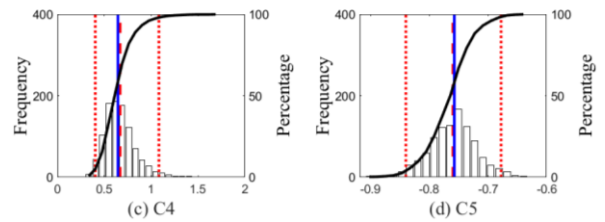
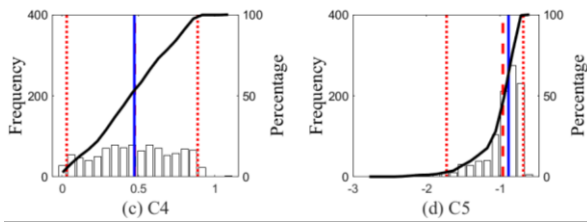
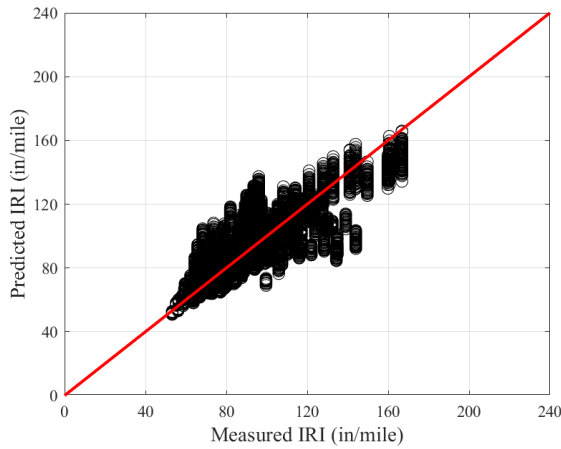


Figure 5-10 Calibration results for transverse cracking using observed data

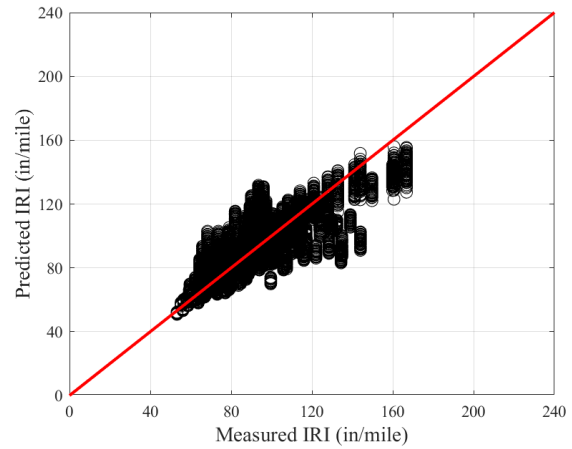
Table 5-14 Summary of calibration results for observed data (IRI - Rigid)

Calibration method	Distribution	SEE	Bias	C_1	C_2	C_3	C_4	NLL	AIC	BIC
Bootstrapping	Global	19.721	11.696	0.820	0.442	1.493	25.24	2094	4191	4199
	Normal	10.208	0.000	0.02	2.194	1.612	24.138	2464	4932	4941
	Exponential	17.503	0.000	1.389	3.953	0.915	8.208	2554	5110	5114
	Gamma	17.304	0.001	1.515	2.518	1.171	7.044	1974	3953	3961
	Log normal	17.772	0.042	1.604	2.546	1.114	7.301	1966	3936	3945
	Negative binomial	10.150	-0.312	0.001	2.229	1.471	27.041	1714	3432	3441
Repeated split sampling (Calibration set)	Normal	10.570	0.000	0.225	2.136	1.510	23.741	1412	2829	2837
	Exponential	18.108	0.000	1.478	3.832	0.893	7.769	1792	3587	3590
	Gamma	17.261	0.000	1.503	2.462	1.201	6.789	1386	2776	2784
	Log normal	17.695	0.038	1.576	2.391	1.176	6.757	1381	2766	2773
	Negative binomial	10.207	-0.316	0.001	2.227	1.476	26.834	1204	2412	2420
Repeated split sampling (Validation set)	Normal	10.654	-0.008	0.225	2.136	1.510	23.741	600	1205	1211
	Exponential	18.208	0.006	1.478	3.832	0.893	7.769	762	1526	1529
	Gamma	17.542	0.098	1.503	2.462	1.201	6.789	590	1185	1191
	Log normal	17.825	0.080	1.576	2.391	1.176	6.757	589	1181	1187
	Negative binomial	10.394	-0.320	0.001	2.227	1.476	26.834	515	1033	1039

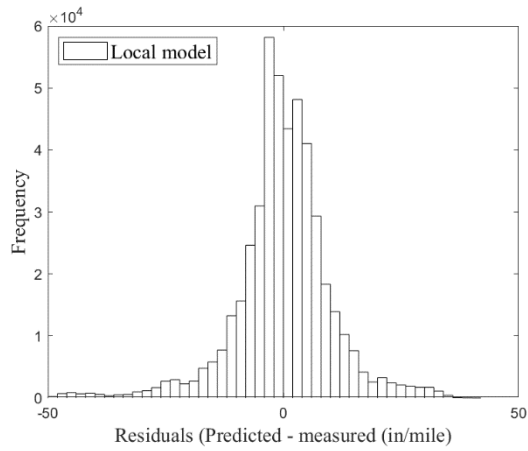
Figure 5-11 shows the calibration results for IRI (using bootstrapping) for the least squares and MLE using a negative binomial distribution. The SEE and bias values for the least squares are 10.208 and 0.000, whereas, for the MLE using negative binomial, they are 10.150 and -0.312, using bootstrap resampling. The predicted vs. measured IRI and distribution of residuals are similar for both methods. Figures 5-11 (e) and (f) show the SEE, bias, and IRI transfer function coefficients distribution for 1000 bootstrap resamples. The least squares method shows lower bias, which can be observed from the distribution of parameters in Figure 5-11. A similar trend is observed in the validation results (Table 5-14), where the least squares method shows better parameter estimates in terms of SEE and bias than the MLE method.



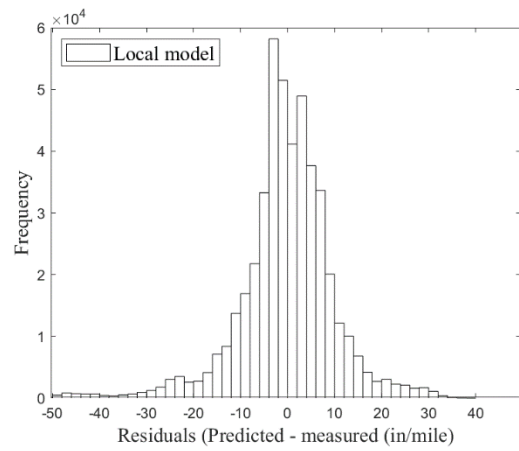
(a) Predicted vs. measured IRI (least squares)



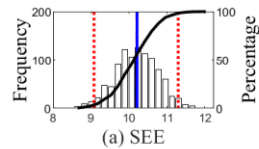
(b) Predicted vs. measured IRI (MLE)



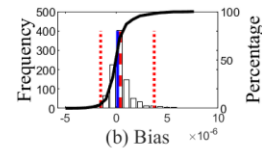
(c) Distribution of residuals (least squares)



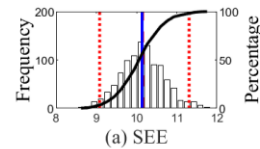
(d) Distribution of residuals (MLE)



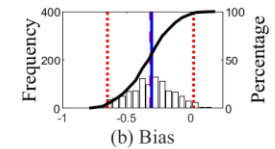
(a) SEE



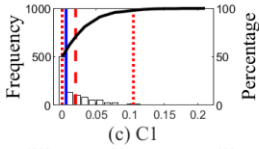
(b) Bias $\times 10^{-6}$



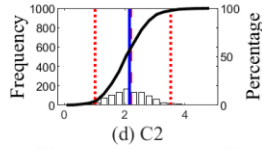
(a) SEE



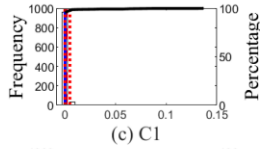
(b) Bias



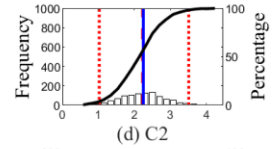
(c) C1



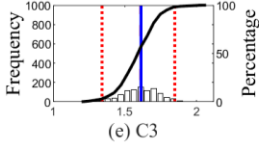
(d) C2



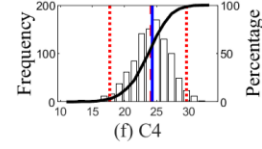
(c) C1



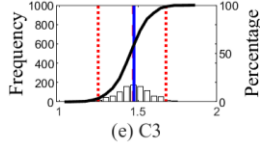
(d) C2



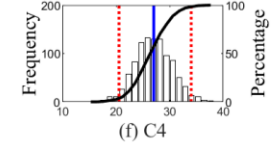
(e) C3



(f) C4



(e) C3



(f) C4

(e) Distribution of parameters (least squares)

(f) Distribution of parameters (MLE)

Figure 5-11 Calibration results for rigid IRI using observed data

5.3 IMPACT OF CALIBRATION ON PAVEMENT DESIGN

Forty-four pavement sections, each for flexible and rigid sections, were designed to assess the impact of calibration on pavement design. It is important to note that the locally calibrated coefficients and standard error equations used for these designs were obtained using the least squares method. The standard error equations are summarized in Chapter 6. Table 5-15 shows the average design thickness for the 44 flexible and rigid sections. These are the final thicknesses based on the following criteria:

- The minimum thickness should be 6.5" for flexible, 9" for JPCP freeway, and 8" for JPCP non-freeway sections.
- A maximum difference of ± 1 inch from the AASHTO93 minimum thickness limits.
- JPCP widened slab sections were designed as standard width (12 feet), and design thicknesses were reduced by a maximum of 1 inch depending on whether the previous conditions were met. This practice is followed because the slab width is a sensitive parameter in the Pavement-ME, giving impractical (very thin) designs.
- The design trials were stopped when a pavement reached a maximum thickness of 16". Few designs fail at even 16", but further increasing the thickness leads to impractical designs. This occurs because a particular design may have inputs (material, traffic, climate) that are not well represented in the global (or local) dataset. Therefore, the Pavement-ME has difficulty providing a practical design outcome. These designs may require changes in the Pavement-ME inputs, and simply changing the thickness cannot achieve a passing design. Furthermore, MDOT is limited by design changes (construction, materials, budget, and design procedures). Therefore, changing the inputs may not be practical.

Table 5-15 Summary of final pavement design thicknesses

Pavement type	Design method	Design thickness (in)		
		Average	Standard deviation	CoV
Flexible	AASHTO93	9.17	2.20	24%
	Pavement-ME previous model	8.86	1.78	20%
	Pavement-ME new calibrated model	8.95	2.27	25%
Rigid	AASHTO93	10.07	1.67	17%
	Pavement-ME global model	9.83	1.63	17%
	Pavement-ME new calibrated model	9.63	1.44	15%

The average design thickness using the newly calibrated models is closer to the AASHTO93 design than the previous model calibration, with an average thickness reduction of 0.22 inches for flexible sections. The average PCC thickness using the new calibrated model is 0.44 inches lower than the AASHTO93 design thickness. Interestingly, for designs using the global model, five sections reached the design thickness of 16 inches, and another five sections reached the design thickness of 6 inches. However, for the design using the locally calibrated model, only one section has a design thickness of 16 inches. Figure 5-12 shows the new calibrated model vs. AASHTO93 design thicknesses. Overall, the average design thickness using the locally calibrated models is slightly lower than the AASHTO93 design thickness for both flexible and rigid sections.

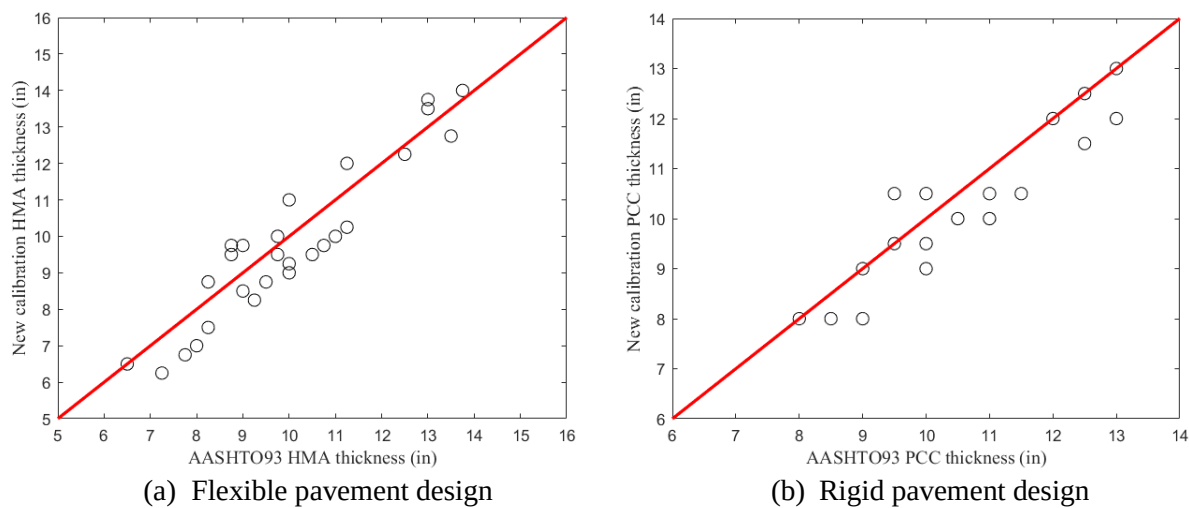
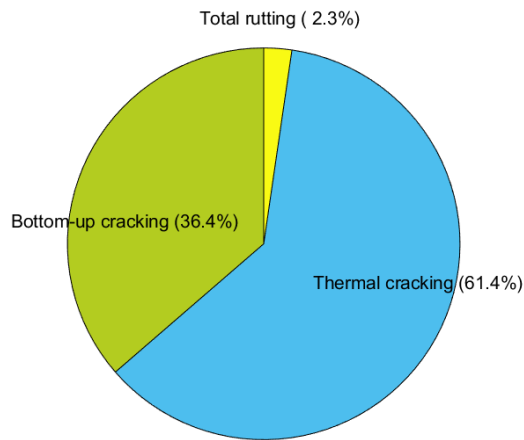
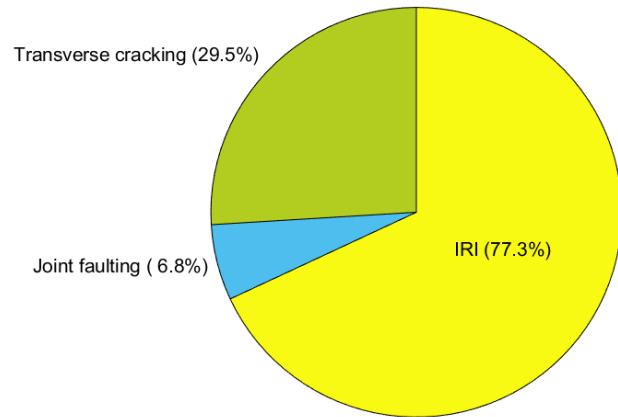


Figure 5-12 New calibrated model vs. AASHTO93 final design thickness

The Pavement-ME designs are based on several distresses, but it is crucial to identify the controlling distress. Figure 5-13 shows the contribution of different controlling distresses. The values shown in Figure 5-13 are the percentage of sections (out of 44) having that critical distress. It should be noted that some sections may have more than one controlling distress. Bottom-up and thermal cracking are the controlling distresses for flexible pavements, whereas transverse cracking and IRI are for rigid pavements. Figure 5-14 compares reliability for critical distress in flexible and rigid pavements. The standard deviation for the previously calibrated model is higher than the newly calibrated model for both bottom-up cracking and thermal cracking in flexible pavements. Also, the standard deviation for the newly calibrated model is higher than the global model for transverse cracking in rigid pavements.

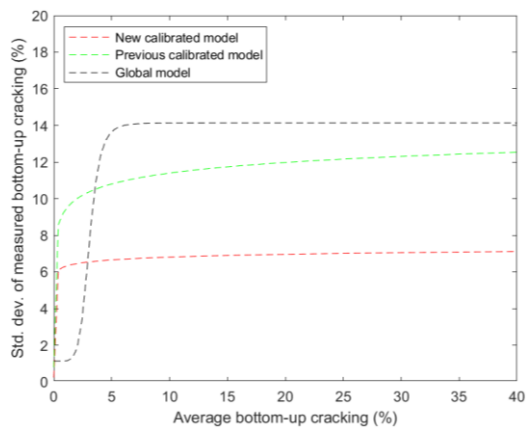


(a) Flexible sections

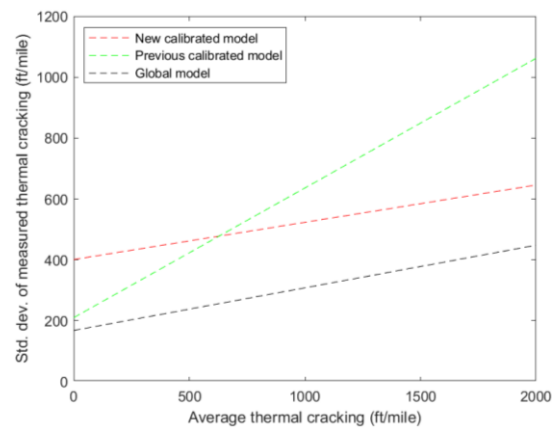


(b) Rigid sections

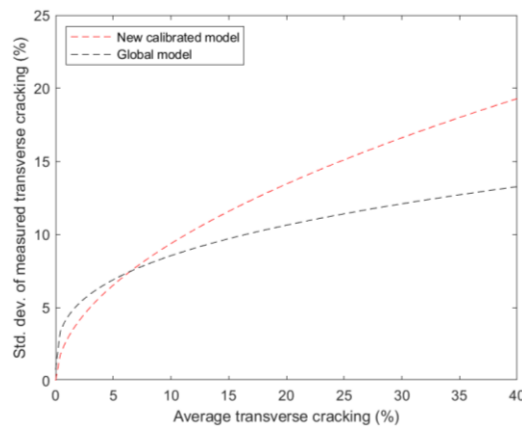
Figure 5-13 Critical distresses for pavement design



(a) Bottom-up cracking (Flexible)



(b) Thermal cracking (Flexible)



(c) Transverse cracking (Rigid)

Figure 5-14 Comparison of reliability for critical distresses

A higher standard deviation in predicted performance is expected to produce a thicker design, but the design results (Table 5-15) show that models with higher standard deviation have lower

design thicknesses. Therefore, these trends indicate that the difference in design thicknesses can be attributed to the calibration coefficients rather than the reliability of these models.

5.4 SENSITIVITY ANALYSIS OF PAVEMENT-ME MODEL COEFFICIENTS

The sensitivity of Pavement-ME model coefficients was estimated using the NSI and SSC methods, as explained in Chapter 4. For NSI calculation, each section was initially run at the global values of transfer function coefficients at 50% reliability. Afterward, each coefficient (one at a time) was varied by -50%, -20%, 20%, and 50%, respectively, from the global values. The change in performance prediction was evaluated for differences in transfer function coefficients to calculate the NSI values. Table 5-16 shows the NSI values for each section in this study and the NSI values from Kim et al. (2014) (26). The NSI values vary significantly among different sections and from Kim et al. (2014). These differences are attributed to the material and climate, ultimately affecting the predicted performance. For example, coefficient C_4 in the IRI model for rigid pavements ranges from 0.06 to 0.23. These values correspond to the coefficient categorized as non-sensitive and sensitive, respectively (60). Similarly, C_2 in bottom-up cracking ranges from -1.3 to -369.5, corresponding to very sensitive and hypersensitive categories. It is important to note that the magnitude of bottom-up cracking in flexible pavements and transverse cracking in rigid pavements was extremely low (close to zero). This has resulted in very high NSI values for C_1 in bottom-up cracking and C_5 in transverse cracking. These values are also significantly different from the ones in Kim et al. (2014). This is mainly because of the difference in magnitude of bottom-up and transverse cracking between the two studies.

The SSCs were calculated and plotted using MATLAB codes using one coefficient at a time and considering other coefficients as constant. A wide range of independent variables have been used since calculating SSCs is a forward problem without data. Figures 5-15 and 5-16 show the SSCs for flexible and rigid pavements. Transfer functions with multiple independent variables have all independent variables shown in the same plot on the x-axis.

Table 5-16 Summary of NSI values for transfer function coefficients

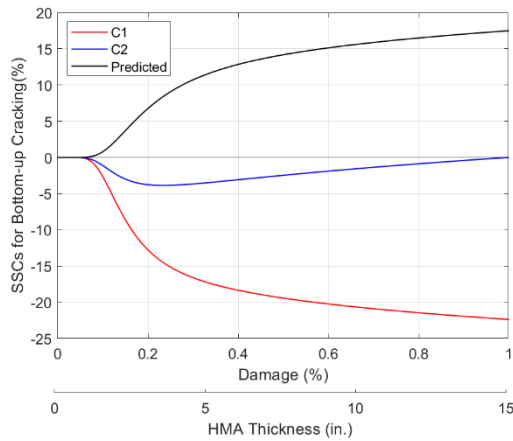
Performance model		Section no.										Average NSI	Kim et al. (26)
		1	2	3	4	5	6	7	8	9	10		
Bottom-up cracking (flexible)	C_1	-61.3	-66.0	-58.5	-88.4	-71	-72.2	-57.1	-58.4	-61.8	-867	-146.2	-11.3
	C_2	-2.9	-5.28	-40.8	-14.14	-1.3	-17.9	-369	-15.1	-94.9	-35.4	-59.75	-2.29
Top-down cracking (flexible)	C_1	-0.59	-0.63	-0.59	0.00	-0.6	-0.67	-0.59	-0.59	-0.72	-0.78	-0.58	NA
	C_2	-2.42	-2.80	-2.37	0.00	-2.7	-3.39	-2.36	-2.42	-4.11	-5.01	-2.76	NA
	C_3	-0.03	-0.18	-0.01	0.00	-0.1	-0.64	0.00	-0.02	0.00	0.00	-0.10	NA
Total rutting (flexible)	β_{1r}	0.23	0.21	0.18	0.13	0.15	0.13	0.19	0.19	0.17	0.18	0.18	1
	β_{s1}	0.24	0.27	0.20	0.28	0.23	0.29	0.17	0.13	0.26	0.29	0.24	1
	β_{sg1}	0.53	0.52	0.62	0.60	0.62	0.58	0.64	0.68	0.57	0.53	0.59	1
IRI (Flexible)	C_1	0.09	0.11	0.07	0.08	0.11	0.09	0.01	0.11	0.06	0.07	0.08	0.15
	C_2	0.02	0.03	0.02	0.00	0.02	0.02	0.02	0.03	0.01	0.01	0.02	0.00
	C_3	0.13	0.13	0.13	0.13	0.06	0.00	0.00	0.00	0.00	0.00	0.06	0.00
	C_4	0.26	0.30	0.22	0.21	0.30	0.25	0.32	0.30	0.25	0.28	0.27	0.31
Transverse cracking (rigid)	C_4	-2.38	-2.38	-2.38	-2.38	-2.38	-2.38	-2.38	-2.38	-2.37	-2.38	-2.38	-0.08
	C_5	-4E4	-1E5	-1E5	-1E6	-1E6	-8E3	-2E3	-3E3	-1E2	-4E4	-2E5	0.20
IRI (rigid)	C_1	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.43
	C_2	0.01	0.00	0.00	0.01	0.01	0.01	0.00	0.00	0.00	0.01	0.01	0.02
	C_3	0.06	0.37	0.34	0.08	0.38	0.25	0.40	0.47	0.48	0.18	0.30	0.07
	C_4	0.15	0.09	0.11	0.23	0.12	0.13	0.07	0.10	0.06	0.13	0.12	0.48

Figures 5-15 and 5-16 show the following observations:

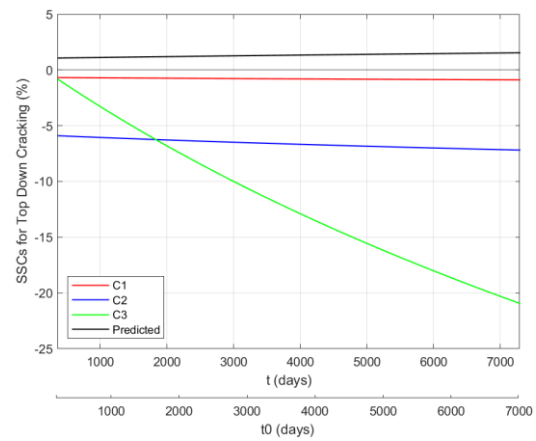
- *Bottom-up cracking (flexible)*: C_1 is more sensitive than C_2 , and C_1 and C_2 are not correlated. Moreover, both C_1 and C_2 are large enough to be confidently estimated. Coefficients with negative SSCs indicate that an increase in the coefficient will decrease predicted performance. Therefore, an increase in C_1 or C_2 will reduce bottom-up cracking.
- *Top-down cracking (flexible)*: The sensitivity of coefficients changes with the independent variables, which are t (analysis time in days) and t_0 (time to crack initiation). Overall, C_3 is the most sensitive coefficient, followed by C_2 . C_1 is the least sensitive coefficient. C_1 and C_2 are correlated, which signifies that only one of them can be estimated with confidence. All coefficients are estimable based on the magnitude of SSCs, and an increase in any of the coefficients will reduce the predicted top-down cracking.
- *Total rutting (flexible)*: Total rutting is a linear model between the individual layer rutting. Subgrade rutting coefficient (β_{sg1}) is the most sensitive, followed by the AC rutting coefficient (β_{1r}). The base rutting coefficient (β_{s1}) is the least sensitive. SSCs for all coefficients are large enough to be estimable and positive.
- *IRI (flexible)*: IRI is a linear relationship between IRI at the time of construction (initial IRI) and other distress (cracking, rutting, etc.). The site factor coefficient is the most sensitive, followed by the total rutting coefficient. The thermal cracking coefficient is the

next sensitive coefficient, while the fatigue cracking coefficient is the least sensitive. All coefficients have positive values for SSCs.

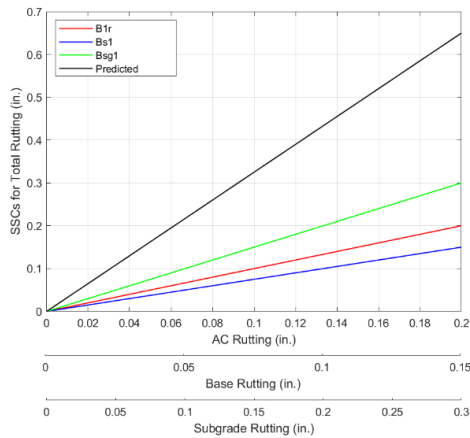
- *Transverse cracking (rigid)*: C_5 is more sensitive than C_4 , and the change in sensitivity with damage can be clearly observed. C_4 and C_5 are not correlated, and the SSCs for both coefficients are large enough to be estimated with confidence.
- *IRI (rigid)*: The transverse cracking coefficient is the most sensitive, and the joint spalling coefficient is the least sensitive. Moreover, the magnitude of SSCs for joint spalling is very low, indicating that the coefficient cannot be estimated with high confidence.



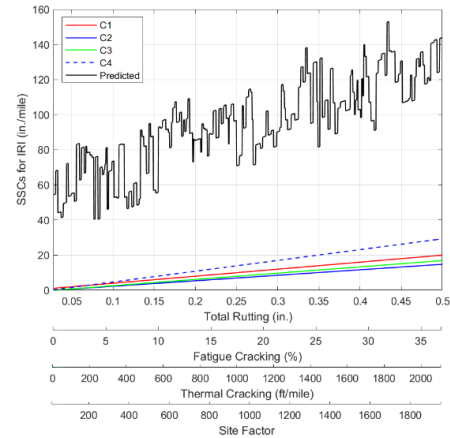
(a) Bottom-up cracking



(b) Top-down cracking



(c) Total rutting



(d) IRI

Figure 5-15 SSCs for flexible pavements

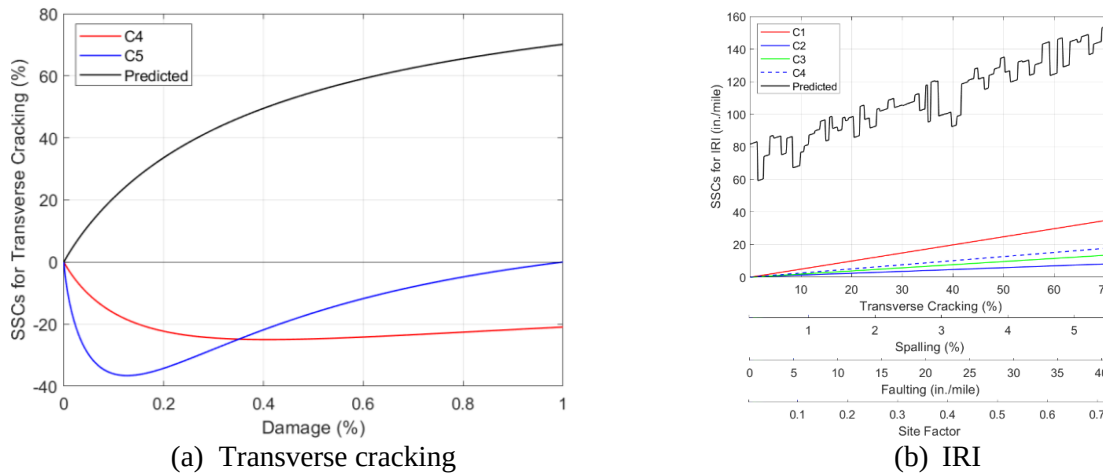


Figure 5-16 SSCs for rigid pavements

SSCs are highly suitable for showing sensitivity for any range of independent variables. For example, the SSC plot for IRI in Figure 5-16 shows that C_1 is the most sensitive coefficient, whereas the NSI values are calculated to show that it is the least sensitive input. This is because of the low values of transverse cracking used to calculate the NSI values. Figure 5-17 shows the SSC plot for IRI in rigid pavements using low values for transverse cracking. It can be observed that at this range of transverse cracking, C_1 is the least sensitive coefficient.

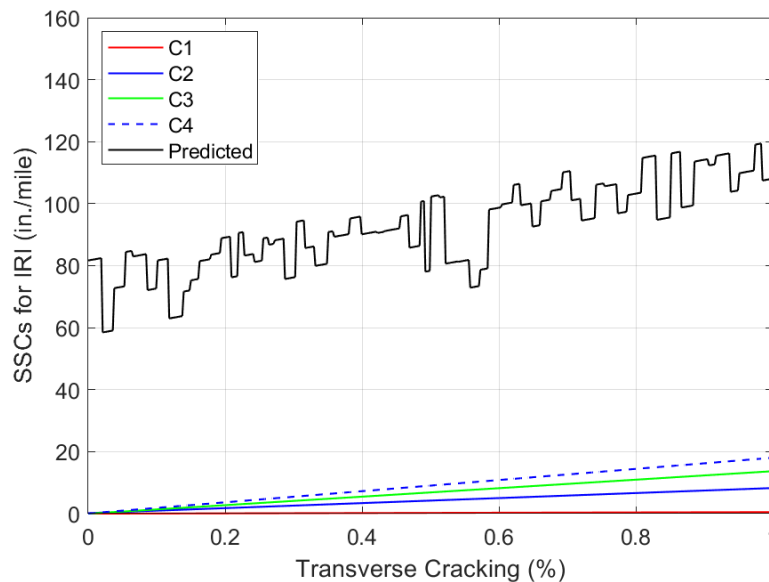


Figure 5-17 SSCs for IRI on small values of transverse cracking

The SSC plot is used to visualize the error in parameter estimation. Moreover, the larger the SSC magnitude, the more confidence in parameter estimation. Calibrating the transverse cracking model in rigid pavements is an example of verification. From Figure 5-16a, C_5 should have less

estimation error than C_4 . Error in estimation for any parameter refers to the relative error, i.e., the ratio of standard error and the parameter value. C_4 and C_5 are not correlated, and the SSCs for both coefficients are large enough to be estimated with confidence. The relative error should be less than 60%; otherwise, the confidence interval of the parameter likely includes zero. In other words, the parameter is not estimable or not statistically different than zero.

The selected rigid pavements were used to calibrate the transverse cracking model and validate the applicability of SSCs. The measured performance data is obtained from the PMS records, and the Pavement-ME inputs are obtained from construction records, material testing results, and the Job Mix Formula (JMF). Figure 5-18 shows the predicted vs. measured transverse cracking for global and locally calibrated model coefficients. Table 5-17 summarizes the standard error of estimate (SEE), bias, and relative error. The local calibration significantly improved the model predictions. Moreover, the relative error for C_5 is less than that for C_4 , with both values less than 60%. The relative error values verify the results from the SSC plot, and therefore, both coefficients can be estimated with confidence.

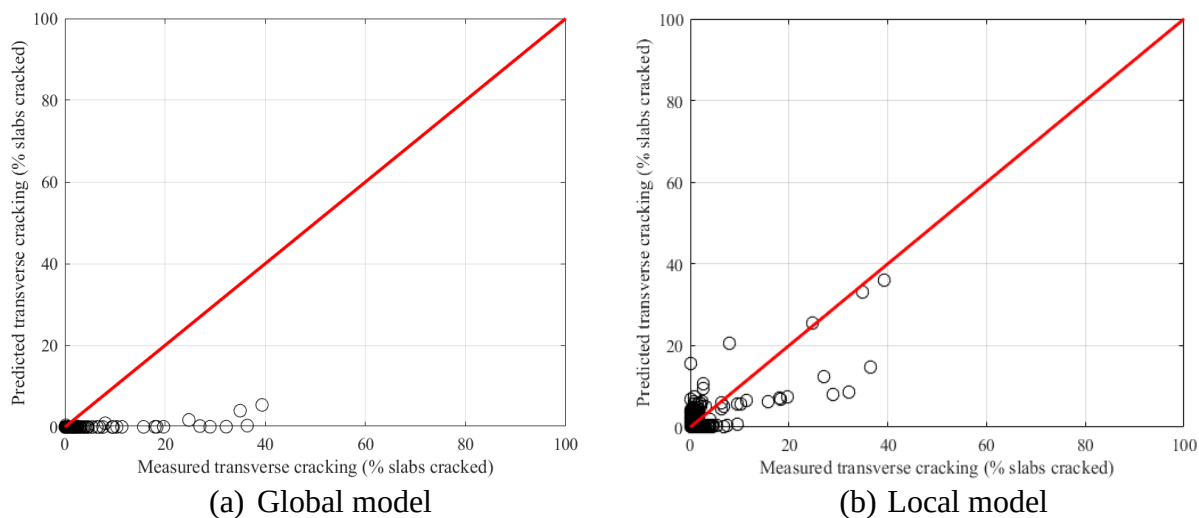


Figure 5-18 Predicted vs. measured transverse cracking in rigid pavements

Table 5-17 Summary of transverse cracking model calibration

Coefficient	Global model			Local model			Relative standard error
	Value	SEE	Bias	Value	SEE	Bias	
C_4	0.52	5.99	-2.39	0.426	3.95	-0.40	20.73%
C_5	-2.17			-0.953			6.14%

5.5 CHAPTER SUMMARY

This chapter summarizes the calibration results for the flexible and rigid Pavement-ME models. Using synthetic and observed data, local calibration was performed using the least squares and MLE methods. Synthetic data was generated using an exponential distribution for bottom-up cracking in flexible pavements and transverse cracking in rigid pavements. MLE results outperformed the least squares method for both sets of synthetic data. Calibration results using observed data showed that MLE provides better parameter estimates for non-normally distributed data. For normally distributed data, MLE and least squares results were comparable. Forty-four sections each for new flexible and rigid pavements were designed using least squares calibration results to assess the impact of calibration on the pavement design. On average, the surface thicknesses using locally calibrated coefficients were thinner than the AASHTO93 design by 0.22 and 0.44 inches for flexible and rigid pavements, respectively. Critical design distresses for flexible pavements include bottom-up and thermal cracking. On the other hand, transverse cracking and IRI control the designs for rigid sections. NSI and SSC methods were used to evaluate the sensitivity of the Pavement-ME transfer function coefficients. Ten sections each, from flexible and rigid pavements, were used to calculate the NSI values and compared with the literature. Results show that SSCs provide a more reliable sensitivity on a range of independent variables rather than a point estimate, unlike NSI. NSI values showed variability among different sections, depending on the magnitude of predicted performance.

CHAPTER 6 - CONCLUSIONS, RECOMMENDATIONS AND FUTURE SCOPE

6.1 KEY FINDINGS

This study introduces a novel calibration approach for the Pavement-ME transfer functions using MLE and compares it with the least squares method. The calibration was performed using synthetic and observed field data. The impact of calibration on pavement design was also assessed. Moreover, the sensitivity of the Pavement-ME model coefficients was also evaluated using the traditional NSI and the SSC approach. The following conclusions can be drawn based on the results.

- The synthetic and observed data distribution for bottom-up cracking in flexible pavements shows skewness, with most data points below 5%. Fitting different distributions over data shows that bottom-up cracking is non-normally distributed. The distribution of observed data for total rutting and IRI shows slight skewness. Moreover, the distribution is close to normal, especially for IRI.
- Calibration results from synthetic data indicate that MLE outperforms the least squares method based on statistical parameters and computational efficiency for flexible pavements. The gamma distribution is the most optimum distribution for MLE, consistently showing SEE and bias values close to zero for the synthetic bottom-up cracking data. The SEE value reduced for MLE results from 3.3 to 0.0 for bootstrapping and 4.4 to 0.0 for repeated split sampling (validation) compared to the least squares results for the dataset with no variability. The dominance of MLE calibration is more evident for datasets with 50% variability, especially in the case of validation.
- For the observed data, the gamma distribution is most suitable for bottom-up cracking and total rutting models, whereas the negative binomial is for the IRI model. The predicted vs. measured plots show less scatter for MLE results than the least squares results for all models. The applicability of MLE is more evident for the bottom-up cracking model. The residual distribution is normally distributed and closer to zero. Moreover, the distribution of parameters is close to a normal distribution, and the bias value is consistently zero, showing the robustness of the calibration results. Calibration of the total rutting model using MLE showed a slight improvement compared to the least squares method, whereas IRI calibration

results for MLE and least squares methods were comparable. This indicates that MLE is more effective for non-normally distributed data.

- The use of MLE on synthetic data for rigid pavements also showed better computational efficiency and applicability to bias-variance tradeoffs compared to the least squares method. The gamma distribution is most suitable for the generated synthetic data for transverse cracking. The mean bias (using bootstrapping) for MLE using gamma distribution is zero for data without and with 50% variability. The SEE values for the least squares method and MLE using gamma distribution are comparable with slightly lower values for MLE. A similar trend is observed in validation results.
- In rigid pavements, the gamma distribution is most suitable for transverse cracking using observed data. The mean bias is consistently near zero using the MLE method for transverse cracking. Calibration using MLE significantly reduces the model bias while keeping the SEE comparable (slightly lower) than the least squares method. Calibration results for IRI using the least squares method and MLE are similar, with the least squares method being somewhat better regarding model bias. The negative binomial is the most suitable distribution for the MLE method.
- The MLE method is proven most effective for skewed and non-normally distributed data, such as bottom-up cracking in flexible pavements and transverse cracking in rigid pavements. In contrast, the least squares method suits data close to a normal distribution, such as IRI. Prior knowledge of distribution is required for the use of MLE.
- Calibration significantly improved performance predictions for both least squares and MLE methods. Resampling methods provide better calibration results with lower SEE and bias and can improve the overall robustness of the MLE approach.
- The average design thicknesses using new calibration coefficients were close to AASHTO93 design thicknesses with a reduction of 0.22 and 0.44 inches in flexible and rigid pavements, respectively. The design thickness using new calibration coefficients was less than the AASHTO93 design thickness for 21 sections and equal for 12 of 44 flexible sections. Similarly, the design thickness using new calibration coefficients was less than the AASHTO93 design thickness for 17 and equal for 23 of 44 rigid sections.
- Thermal cracking is the most critical distress for flexible sections, with 61.4%, followed by bottom-up cracking, with a 36.4% contribution. The contribution of total rutting was 2.3%.

None of the sections had top-down cracking or IRI as their critical distress. In rigid sections, IRI controlled distress with 77.3%, followed by transverse cracking with a 29.5% contribution. Joint faulting had the most negligible contribution of 6.8%. Comparison between the standard deviation of different models indicates that the differences in design thicknesses come from the calibration coefficients rather than the reliability for both flexible and rigid sections.

- The sensitivity analysis showed that NSI values differed for each section in both flexible and rigid pavements. These sections have been designed using different Pavement-ME inputs, resulting in a wide range of performance predictions and, ultimately, a range of NSI values. The bottom-up cracking predictions in flexible and transverse cracking predictions in rigid sections were extremely low (close to zero). This resulted in very high NSI values, which are unreliable. The coefficient C_1 for IRI in rigid sections is also zero because of the low magnitude of transverse cracking. NSI values are variable and depend on the magnitude of predicted distresses. Moreover, the Pavement-ME inputs (material, traffic, and climatic) are required for NSI calculations.
- SSCs provide a convenient visual representation of the sensitivity of different transfer function coefficients over a continuous range of independent variables, unlike NSI, which is a point estimate. SSCs for transverse cracking and IRI for rigid sections show that the sensitivity changes at different ranges of the independent variable. It also indicates any correlations between different coefficients and confidence in estimation. Calculation of SSCs is a forward problem and does not require any input data. Therefore, a user only needs a mathematical model (the transfer functions) and can calculate SSCs on any range of independent variables.
- NSI and SSCs provide a measure of sensitivity, but it is convenient to rank transfer function coefficients for straightforward interpretation. Table 6-1 shows the ranking of transfer function coefficients based on different methods. The order using SSCs is based on the overall sensitivity in the entire range of independent variables. As previously shown, this sensitivity might change for a limited range of independent variables. Coefficients with the same NSI values have been ranked the same. For example, all rutting coefficients in Kim et al. (2014) (26) study have been ranked 1 as they all have the same NSI values. Some models (e.g., bottom-up cracking and transverse cracking) have similar rankings using different

methods, whereas others (e.g., IRI for rigid pavements) have significant differences. These differences make it challenging to estimate the most sensitive coefficients truly. Therefore, SSCs can help obtain a continuous range of sensitivity rather than a point estimate.

Table 6-1 Rank of transfer function coefficients based on different methods

Pavement type	Performance model	Coefficient	NSI	SSCs	Kim et al. (2014) (26)
Flexible	Bottom-up cracking	C_1	1	1	1
		C_2	2	2	2
	Top-down cracking	C_1	2	3	NA
		C_2	1	2	NA
		C_3	3	1	NA
	Total rutting	β_{1r}	3	2	1
		β_{s1}	2	3	1
		β_{sg1}	1	1	1
	IRI	C_1	3	2	2
		C_2	4	4	3
		C_3	2	3	3
		C_4	1	1	1
Rigid	Transverse cracking	C_4	2	2	2
		C_5	1	1	1
	IRI	C_1	4	1	2
		C_2	3	4	4
		C_3	1	3	3
		C_4	2	2	1

6.2 RECOMMENDED CALIBRATION COEFFICIENTS

Tables 6-2 and 6-3 summarize the recommended calibration coefficients and reliability equations for flexible and rigid pavements. These results were obtained using the least squares method and validated with extensive pavement designs. The detailed results of pavement designs are shown in Chapter 5.

Table 6-2 Flexible pavement recommended calibration coefficients and standard error equations

Performance prediction model		Local coefficient	Standard error
Bottom-up cracking (Option a)		$C_1 = 0.2320$ $C_2 = 0.6998 \text{ (hac} < 5 \text{ in)}$ $C_2 = (0.867 + 0.2583 * hac) * 0.2204 \text{ (5 in} \leq \text{hac} \leq 12 \text{ in)}$ $C_2 = 0.8742 \text{ (hac} > 12 \text{ in)}$	$S_{e(BU)} = 0.2262 + \frac{14.2349}{1 + \exp(0.2958 - 0.1441 \log(Crack))}$
Bottom-up cracking (Option b)		$C_1 = 0.2540$ $C_2 = 0.7303 \text{ (hac} < 5 \text{ in)}$ $C_2 = (0.867 + 0.2583 * hac) * 0.2692 \text{ (5 in} \leq \text{hac} \leq 12 \text{ in)}$ $C_2 = 1.0678 \text{ (hac} > 12 \text{ in)}$	$S_{e(BU)} = 4.4396 + \frac{25.4391}{1 + \exp(4.3119 - 2.2778 \log(Crack))}$
Top-down cracking		$K_{L1} = 64271618$ $K_{L2} = 0.90$ $K_{L3} = 0.09$ $K_{L4} = 0.101$ $K_{L5} = 3.260$ $C_1 = 0.30$ $C_2 = 1.155$ $C_3 = 1$	$S_{e(TD)} = 0.6417 \times TOP + 0.5014$
Rutting	HMA	$\beta_{1r} = 0.148$ $\beta_{2r} = 0.7$ $\beta_{3r} = 1.3$	$S_{e(HMA)} = 0.1481(RUT_{HMA})^{0.4175}$
	Base/subgrade	$\beta_{s1} = 0.301$ $\beta_{sg1} = 0.070$	$S_{e(base)} = 0.0411(RUT_{base})^{0.3656}$ $S_{e(subgrade)} = 0.0728(RUT_{subgrade})^{0.5456}$
Thermal cracking		$K = 0.85$	$S_{e(TC)} = 0.1223(TC) + 400.9$
IRI		$C_1 = 42.874, C_2 = 0.102$ $C_3 = 0.0081, C_4 = 0.003$	Internally determined by the software

Table 6-3 Rigid pavement recommended calibration coefficients and standard error equations

Performance prediction model		Local coefficient	Standard error
Transverse cracking		$C_4 = 0.415$ $C_5 = -0.965$	$S_{e(CRK)} = 2.9004(CRK)^{0.5074}$
Transverse joint faulting		$C_1 = 0.6$ $C_2 = 1.611$ $C_3 = 0.00217$ $C_4 = 0.00444$ $C_5 = 250$ $C_6 = 0.2$ $C_7 = 7.3$ $C_8 = 400$	$S_{e(Fault)} = 0.0919(Fault)^{0.2249}$
IRI		$C_1 = 0.0942$ $C_2 = 1.5471$ $C_3 = 1.7970$ $C_4 = 23.7529$	Internally determined by the software

6.3 PRACTICAL IMPLICATIONS

This study provides a framework for the local calibration of performance models. Highway agencies can leverage the results for better design and adaptation of the Pavement-ME for local conditions. The critical implications include:

- The recorded performance data may have irregularities due to measurement errors and limitations in distress identification. Moreover, the recorded performance data may require conversion to the Pavement-ME units, which involves several assumptions. It may cause anomalies in the measured performance data and may not be practical to use directly for calibration. Therefore, analyzing the raw performance data and filtering it (if required) is recommended for practicality.
- It is worth mentioning that the calibration process and pavement design were simultaneously executed. For every set of calibration coefficients, pavements were designed, and the calibration was improved based on the results. Pavement design is one of the most crucial calibration process steps and is often not considered in practice. It is recommended that the calibration results should not be based only on statistical parameters (SEE, bias, etc.) but also on practical engineering judgments.
- Identifying critical design distress types is crucial. By understanding which distress types are most relevant to their region, agencies can develop mitigation and maintenance strategies leading to longer pavement service lives. For example, thermal cracking is critical in Michigan for flexible pavements. MDOT can mitigate the occurrence of cracking by using modified and improved binders.
- It is recommended that local calibrations be performed every six years when more time series data points (e.g., three data points in Michigan) are available for the already selected and new pavement sections.

SSCs can help agencies improve their local calibration process. The advantages and interpretation of the SSC plots are described in Chapters 4 and 5. The following approach is recommended to leverage these SSC plots before starting the local calibration process:

- Run Pavement-ME to identify the magnitude of independent variables for each model. For example, one should know the range of damage values for transverse cracking in rigid pavements.
- Obtain the sensitivity of each calibration coefficient from the SSC plots for the respective range of independent variables.
- Ensure that the SSCs for each coefficient are large enough (the maximum value of SSC should be at least 10% of the largest value of the dependent variable). For example, the maximum SSC values for C_4 and C_5 are 25% and 38%, respectively, in transverse cracking

in rigid pavements. These SSC values exceeded 10% of the maximum predicted transverse cracking. Moreover, the SSCs should not be correlated (SSCs for different coefficients should not show similar trends).

- If the SSCs are not large enough, one can not estimate those coefficients with sufficient confidence, i.e., they may be insignificant. On the other hand, if coefficients are correlated, both coefficients cannot be simultaneously estimated. For example, coefficients C_1 and C_2 in top-down cracking for flexible pavements show a correlation; therefore, only one should be calibrated. Calibration of C_2 is recommended since the magnitude of SSC for C_2 is higher.
- Ensure that the relative error is lower for the more sensitive coefficients and is not more than 60% for any coefficient.
- The SSCs can highlight the most significant coefficients for a range of independent variables. That can help in diverting more attention to those coefficients during local calibration. For example, in the rigid IRI model, C_1 is the least sensitive for lower transverse cracking (less than 1%), and C_1 is the most sensitive for higher transverse cracking.

6.4 REVIEW OF CAT TOOL

This study used The CAT tool to calibrate the thermal cracking model in flexible and joint faulting models in rigid pavements. CAT provides a convenient alternative for those models where rerunning Pavement-ME is required. The advantages and limitations of the CAT tool are summarized below:

Advantages of CAT

- CAT provides good visualization of the input data and experimental matrix of the *.dgp files. It helps quickly glance at the overall data and identify any outliers or biases.
- It has default validation of the optimized coefficients, which helps to verify the model on an independent set of sections.
- It provides sufficient descriptive statistics for calibration results and a linear model showing the effect of different Pavement-ME inputs on the overall calibration.
- It helps visualize the change in error and bias for each iteration, making it easier to identify local minima in the given range.
- It assists in evaluating the impact of the number of bins on the reliability model.

Limitations of CAT

- Pavement-ME (.dgp) files, once uploaded, cannot be deleted. Also, trials for optimizing calibration coefficients are run; they cannot be deleted or paused. This makes the nomenclature of the .dgp files challenging, and trials should be run meticulously.
- The limit to the total number of combinations of calibration coefficients is 100. Hence, all calibration coefficients cannot be changed simultaneously for several increments.
- Since the number of increments is fixed to 100, the coefficients must be changed systematically by reducing the range provided. Also, not more than three coefficients can be involved in one trial run for a reasonable range and number of increments. It makes the optimization process cumbersome, and some prior experience is required to recalibrate with optimum time and effort.
- The computation time is comparatively large. For example, for 100 pavement sections, changing a total of two calibration coefficients with five increments each makes it a total of $100 \times 5 \times 5 = 2500$ Pavement-ME runs, which takes a computation time of around 29 hours. Therefore, considerable computational time is required, especially when the number of sections is large.
- The same sections cannot be used for different projects using different measured data. Changing the measured data changes it in all existing (already run) projects.

6.5 FUTURE SCOPE OF THIS STUDY

The scope of this study is limited to new flexible and rigid pavements. Moreover, bottom-up cracking, total rutting, IRI models for flexible pavements, and transverse cracking and IRI models for rigid pavements were calibrated using the four distributions mentioned: exponential, gamma, log-normal, and negative binomial. Using an exponential distribution, the MLE methodology was validated using synthetic data for bottom-up cracking in flexible and transverse cracking in rigid pavements. The following can be explored as part of future studies:

- The MLE approach can be extended to calibrate other Pavement-ME models and models for rehabilitated pavements. Different probability distributions can be explored as part of future research.

- This methodology can be validated using synthetic data for different Pavement-ME transfer functions. Moreover, synthetic data can be generated using different distributions and variability.
- Further studies can be conducted to estimate the impact of varying calibration approaches on pavement design.
- Top-down cracking model calibration improved the SEE and bias but did not provide realistic results, i.e., high SEE. Furthermore, the top-down cracking predictions didn't vary for different sections, producing the same predictions. The Pavement-ME limits the thermal cracking prediction to 2112 ft/mile, but the measured data showed several records of thermal cracking above 2112 ft/mile. Also, the thermal cracking coefficient in the current version is changed and is a function of MAAT. This made the calibration of the thermal cracking model challenging. Due to the model's limitations, although the SEE and bias were improved after local calibration, the thermal cracking model still showed high variability. The top-down and thermal cracking models in flexible pavements should be improved, especially considering thermal cracking is critical.
- The SSCs can be used for sensitivity analysis in other Pavement-ME models, apart from the transfer functions.

REFERENCES

1. Pierce LM, McGovern G. Implementation of the AASHTO mechanistic-empirical pavement design guide and software. NCHRP Synthesis 457. 2014.
2. ARA. Guide for mechanistic empirical design of new and rehabilitated pavement structures. ARA, Inc, ERES Consultants Division: Champaign, IL. 2004.
3. AASHTO. Mechanistic–Empirical Pavement Design Guide, Interim Edition: A Manual of Practice. American Association of State Highway and Transportation Officials: Washington, DC. 2008.
4. NCHRP. Guide for mechanistic empirical design of new and rehabilitated pavement structures. NCHRP 1-37A. 2004.
5. AASHTO. Guide for design of pavement structures. 1993.
6. Li Q, Xiao DX, Wang KC, Hall KD, Qiu Y. Mechanistic-empirical pavement design guide (MEPDG): a bird's-eye view. *Journal of Modern Transportation*. 2011;19:114-33.
7. AASHTO. Guide for the Local Calibration of the Mechanistic-Empirical Pavement Design Guide. 2010.
8. Haider SW, Brink WC, Buch N. Local calibration of rigid pavement performance models using resampling methods. *International Journal of Pavement Engineering*. 2017;18(7):645-57.
9. Haider SW, Musunuru G, Buch N, Brink WC. Local Recalibration of JPCP Performance Models and Pavement-ME Implementation Challenges in Michigan. *Journal of Transportation Engineering, Part B: Pavements*. 2020;146(1):04019037.
10. Haider SW, Buch N, Brink W, Chatti K, Baladi G. Preparation for implementation of the mechanistic-empirical pavement design guide in Michigan, part 3: local calibration and validation of the pavement-ME performance models. Michigan. Dept. of Transportation. Office of Research Administration; 2014.
11. Dong S, Zhong J, Tighe SL, Hao P, Pickel D. Approaches for local calibration of mechanistic-empirical pavement design guide joint faulting model: a case study of Ontario. *International Journal of Pavement Engineering*. 2020;21(11):1347-61.
12. Grogg M, Pierce L, Smith K. Mechanistic-Empirical Pavement Design Guide (MEPDG) Implementation Roadmap. 2022.
13. Hall KD, Xiao DX, Wang KC. Calibration of the mechanistic–empirical pavement design guide for flexible pavement design in Arkansas. *Transportation Research Record*. 2011;2226(1):135-41.
14. Tarefder R, Rodriguez-Ruiz JI. Local calibration of MEPDG for flexible pavements in New Mexico. *Journal of Transportation Engineering*. 2013;139(10):981-91.
15. Tabesh M, Sakhaeifar MS. Local calibration and Implementation of AASHTOWARE Pavement ME performance models for Oklahoma pavement systems. *International Journal of Pavement Engineering*. 2021:1-12.
16. Hall KD, Xiao DX, Wang KC. Calibration of the ME Design Guide. Arkansas State Highway and Transportation Department; 2014.

17. Mallela J, Titus-Glover L, Sadasivam S, Bhattacharya B, Darter M, Von Quintus H. Implementation of the AASHTO mechanistic-empirical pavement design guide for Colorado. Colorado. Dept. of Transportation. Research Branch; 2013.
18. Velasquez R, Hoegh K, Yut I, Funk N, Cochran G, Marasteanu M, et al. Implementation of the MEPDG for new and rehabilitated pavement structures for design of concrete and asphalt pavements in Minnesota. 2009.
19. Von Quintus HL, Moulthrop JS. Mechanistic-Empirical Pavement Design Guide Flexible Pavement Performance Prediction Models: Volume I Executive Research Summary. Montana. Dept. of Transportation. Research Programs; 2007.
20. Titus-Glover L, Mallela J. Guidelines for Implementing NCHRP 1-37A ME Design Procedures in Ohio: Volume 4--MEPDG Models Validation & Recalibration. Ohio. Dept. of Transportation; 2009.
21. Williams RC, Shaidur R. Mechanistic-empirical pavement design guide calibration for pavement rehabilitation. Oregon. Dept. of Transportation. Research Section; 2013.
22. Gassman SL, Rahman MM. Calibration of the AASHTO pavement design guide to South Carolina conditions-phase I. South Carolina. Dept. of Transportation; 2016.
23. Darter MI, Titus-Glover L, Von Quintus H. IMPLEMENTATION OF THE MECHANISTIC-EMPIRICAL PAVEMENT DESIGN GUIDE IN UTAH: VALIDATION, CALIBRATION, AND DEVELOPMENT OF THE UDOT MEPDG USER'S GUIDE. Utah Department of Transportation, Research Division; 2009.
24. Li J, Pierce LM, Uhlmeier J. Calibration of flexible pavement in mechanistic-empirical pavement design guide for Washington state. Transportation Research Record. 2009;2095(1):73-83.
25. Darter MI, Von Quintus H, Bhattacharya BB, Mallela J. Calibration and implementation of the AASHTO mechanistic-empirical pavement design guide in Arizona. Arizona. Dept. of Transportation. Research Center; 2014.
26. Kim S, Ceylan H, Ma D, Gopalakrishnan K. Calibration of pavement ME design and mechanistic-empirical pavement design guide performance prediction models for Iowa pavement systems. Journal of Transportation Engineering. 2014;140(10):04014052.
27. Sun X, Han J, Parsons RL, Misra A, Thakur JK. Calibrating the mechanistic-empirical pavement design guide for Kansas. Kansas. Dept. of Transportation. Bureau of Materials & Research; 2015.
28. Haider SW, Kutay ME, Cetin B, Singh RR, Muslim HB, Santos C, et al. Testing Protocol, Data Storage, and Recalibration for Pavement-ME Design. Michigan. Dept. of Transportation. Research Administration; 2023.
29. Kim YR, Jadoun FM, Hou T, Muthadi N. Local calibration of the MEPDG for flexible pavement design. North Carolina State University. ; 2011.
30. Banerjee A, Aguiar-Moya J, Prozzi J. Texas experience using LTPP for calibration of the MEPDG permanent deformation models. Transportation Research Record. 2009;2094:12-20.

31. Bhattacharya BB, Von Quintus HL, Darter MI. Implementation and local calibration of the MEPDG transfer functions in Wyoming. Wyoming. Dept. of Transportation; 2015.
32. Titus-Glover L, Rao C, Sadasivam S. Local Calibration of the Pavement ME for Missouri. Missouri. Department of Transportation. Construction and Materials Division; 2020.
33. Von Quintus HL, Darter MI, Bhattacharya BB, Titus-Glover L. Calibration of the MEPDG transfer functions in Georgia: task order 2 report. Georgia. Dept. of Transportation; 2015.
34. Wu Z, Xiao DX, Zhang Z. Research Implementation of AASHTOWare Pavement ME Design in Louisiana. Transportation Research Record. 2016;2590(1):1-9.
35. Nair H, Saha B, Merine G. Developing an Implementation Strategy for Virginia Department of Transportation Pavement Rehabilitation Design Using Mechanistic-Empirical Concepts. Virginia Transportation Research Council (VTRC); 2022.
36. Huang B, Gong H, Shu X. Local Calibration of Mechanistic-Empirical Pavement Design Guide in Tennessee. Tennessee. Department of Transportation; 2016.
37. Bayomy F, Muftah A, Kassem E, Williams C, Hasnat M. Calibration of the AASHTOWare Pavement ME Design Software for PCC Pavements in Idaho. Idaho. Transportation Department; 2019.
38. Chen L, Zhang F, Zhou C. Maximum Likelihood Estimation of Parameters for Advanced Continuously Reinforced Concrete Pavement (CRCP) Punchout Calibration Model. Advances in Civil Engineering. 2021;2021:1-8.
39. Myung IJ. Tutorial on maximum likelihood estimation. Journal of mathematical Psychology. 2003;47(1):90-100.
40. Dwivedi R, Singh C, Yu B, Wainwright M. Revisiting minimum description length complexity in overparameterized models. Journal of Machine Learning Research. 2023;24(268):1-59.
41. Franco JCG. Maximum Likelihood Estimation of a Mean Reverting Process. 2023.
42. Pan J-X, Fang K-T. Maximum likelihood estimation. Growth curve models and statistical diagnostics. 2002:77-158.
43. Bauke H. Parameter estimation for power-law distributions by maximum likelihood methods. The European Physical Journal B. 2007;58:167-73.
44. Zhang Y, Callan J, editors. Maximum likelihood estimation for filtering thresholds. Proceedings of the 24th annual international ACM SIGIR conference on Research and development in information retrieval; 2001.
45. Rayner GD, MacGillivray HL. Numerical maximum likelihood estimation for the g-and-k and generalized g-and-h distributions. Statistics and Computing. 2002;12(1):57-75.
46. Lio W, Liu B. Uncertain maximum likelihood estimation with application to uncertain regression analysis. Soft Computing. 2020;24:9351-60.
47. Miner MA. Cumulative damage in fatigue. 1945.
48. AASHTO. Top-Down Cracking Enhancement <https://me-design.com/MEDesign/Documents>. 2020.

49. Ling M, Luo X, Chen Y, Gu F, Lytton RL. Mechanistic-empirical models for top-down cracking initiation of asphalt pavements. *International Journal of Pavement Engineering*. 2020;21(4):464–73.
50. Fick SB. Evaluation of the AASHTO empirical and mechanistic-empirical pavement design procedures using the AASHO road test: University of Maryland, College Park; 2010.
51. AASHTO. Mechanistic-empirical pavement design guide: A manual of practice, 2nd Edition. 2015.
52. Wu Z, Xiao DX, Zhang Z, Temple WH. Evaluation of AASHTO Mechanistic-Empirical Pavement Design Guide for designing rigid pavements in louisiana. *International Journal of Pavement Research and Technology*. 2014;7(6):405.
53. Nam Tran P, Robbins MM, Rodezno C. Pavement ME Design–Impact of Local Calibration, Foundation Support, and Design and Reliability Thresholds. 2017.
54. Mu F, Mack JW, Rodden RA. Review of national and state-level calibrations of AASHTOWare Pavement ME design for new jointed plain concrete pavement. *International Journal of Pavement Engineering*. 2018;19(9):825-31.
55. Singh RR, Haider SW, Schenkel JP. Impact of Local Calibration on Pavement Design in Michigan. *Journal of Transportation Engineering, Part B: Pavements*. 2024.
56. Buch N, Chatti K, Haider SW, Manik A. Evaluation of the 1–37A Design Process for New and Rehabilitated JPCP and HMA Pavements. 2008.
57. Harsini I, Brink WC, Haider SW, Chatti K, Buch N, Baladi GY, et al. Sensitivity of Input Variables for Flexible Pavement Rehabilitation Strategies in the MEPDG. 2013.
58. Mallela J, Titus-Glover L, Bhattacharya B, Darter M, Von Quintus H. Idaho AASHTOWare pavement ME design user's guide, version 1.1. Idaho. Transportation Dept.; 2014.
59. Brink WC, Harsini I, Haider SW, Buch N, Chatti K, Baladi GY, et al. Sensitivity of input variables for rigid pavement rehabilitation strategies in the MEPDG. *Airfield and Highway Pavement 2013: Sustainable and Efficient Pavements*. 2013:528-38.
60. Schwartz C, Li R, Kim S, Ceylan H, Gopalakrishnan K. Sensitivity evaluation of MEPDG performance prediction, NCHRP 1–47. 2011.
61. Ceylan H, Gopalakrishnan K, Kim S, Schwartz CW, Li R. Global sensitivity analysis of jointed plain concrete pavement mechanistic–empirical performance predictions. *Transportation research record*. 2013;2367(1):113-22.
62. Hall KD, Beam S. Estimating the sensitivity of design input variables for rigid pavement analysis with a mechanistic–empirical design guide. *Transportation Research Record*. 2005;1919(1):65-73.
63. Dolan KD, Mishra DK. Parameter Estimation in Food Science. *Annual Review of Food Science and Technology*. 2013;4(1):401-22.
64. Beck JV, Arnold KJ. Parameter estimation in engineering and science: James Beck; 1977.

65. Geeraerd A, Valdramidis V, Van Impe J. GInaFiT, a freeware tool to assess non-log-linear microbial survivor curves. *International journal of food microbiology*. 2005;102(1):95-105.
66. Dolan KD. Estimation of kinetic parameters for nonisothermal food processes. *Journal of Food Science*. 2003;68(3):728-41.
67. Mishra DK, Dolan KD, Beck JV, Ozadali F. Use of scaled sensitivity coefficient relations for intrinsic verification of numerical codes and parameter estimation for heat conduction. *Journal of Verification, Validation and Uncertainty Quantification*. 2017;2(3):031005.
68. You Z, Yang X, Hiller J, Watkins D, Dong J. Improvement of Michigan climatic files in pavement ME design. Michigan. Dept. of Transportation; 2015.
69. Hasnat M, Singh RR, Kutay ME, Bryce J, Haider SW, Cetin B. Comparative study of different condition indices using Michigan department of transportation's flexible distress data. *Transportation Research Record*. 2023;2677(8):400-13.
70. Miller JS, Bellinger WY. Distress identification manual for the long-term pavement performance program. United States. Federal Highway Administration.; 2003.
71. Kutay ME, Jamrah A. Preparation for Implementation of the Mechanistic-Empirical Pavement Design Guide in Michigan: Part 1–HMA Mixture Characterization. 2013.
72. MDOT. Michigan DOT User Guide for Mechanistic-Empirical Pavement Design. Michigan Department of Transportation Lansing, MI, USA; 2021.
73. Baladi GY, Thottempudi A, Dawson T. Backcalculation of unbound granular layer moduli. 2011.
74. Baladi G, Dawson T, Sessions C. Pavement subgrade MR design values for michigan's seasonal changes, final report. Michigan Department of Transportation, Construction and Technology Division, PO Box. 2009;30049.
75. Singh RR, Haider SW, Kutay ME, Cetin B, Buch N. Impact of Climatic Data Sources on Pavement Performance Prediction in Michigan. *Journal of Transportation Engineering, Part B: Pavements*. 2022;148(3):04022048.

APPENDIX

This chapter summarizes the results of the Pavement-ME models calibrated using the least squares method only. These models include bottom-up cracking (Option a), top-down cracking, thermal cracking, and rutting (Method 1) models for flexible pavements and joint faulting models for rigid pavements.

BOTTOM-UP CRACKING MODEL (OPTION A)

No Sampling

In no sampling, the entire dataset was used for calibration. The error between the predicted and measured fatigue cracking was minimized. Figure A-1 shows the predicted versus measured bottom-up for the global and locally calibrated models. The global model underpredicts bottom-up cracking. Table A-1 shows the local calibration results. The SEE is reduced from 8.28 to 8.08, whereas the bias is reduced from -4.90 to 0.17. Figure A-2 shows the fatigue damage curve and the measured and locally predicted bottom-up cracking with time. These measured and predicted cracking values are for the same sections and at the same ages. Figure A-2 shows that local predictions are close to the measured values.

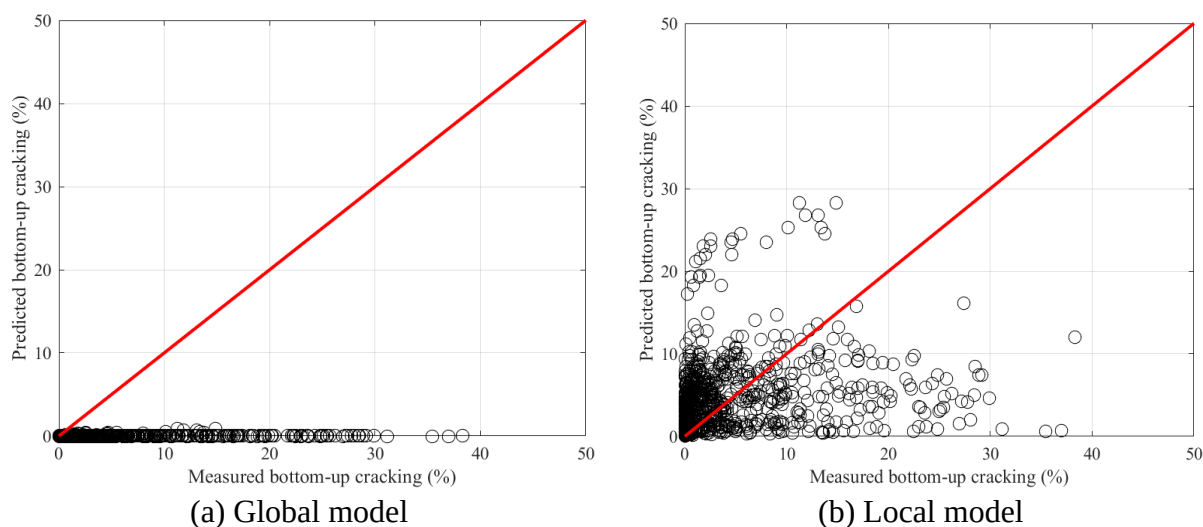
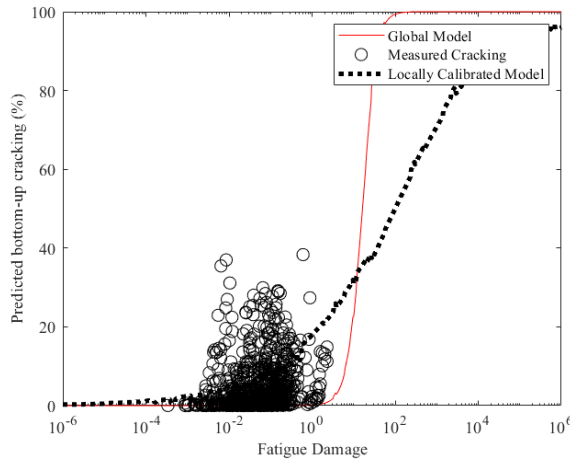
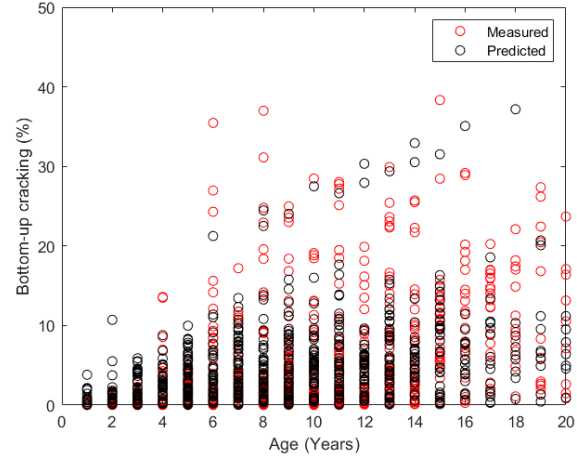


Figure A-1 Predicted vs. measured bottom-up cracking (No sampling)



(a) Fatigue damage



(b) Measured and predicted time series

Figure A-2 Local calibration results for bottom-up cracking (No sampling)

Table A-1 Local calibration summary for bottom-up cracking (No sampling)

Parameter	Global model	Local model
SEE (% total lane area)	8.28	8.08
Bias (% total lane area)	-4.90	0.17
C_1	1.31	0.22
C_2 ($h_{ac} < 5$ in.)	2.1585	0.66
C_2 ($5 \text{ in.} \leq h_{ac} \leq 12 \text{ in.}$)	$(0.867 + 0.2583 * h_{ac}) * 1$	$(0.867 + 0.2583 * h_{ac}) * 0.22$

Split Sampling

Split sampling was used with a random split of 70% sections for the calibration set and the rest 30% for the validation set. Figure A-3 shows the predicted vs. measured bottom-up cracking for the calibration and validation sets. The validation set shows a similar trend as the calibration set. Table A-2 summarizes the local calibration results. Though SEE is higher than the global model, bias is significantly improved from -4.54 to 0.7018 in the validation set. Overall, the validation results are satisfactory.

Table A-2 Local calibration summary for bottom-up cracking (split sampling)

Parameter	Global model	Local model	Validation
SEE (% total lane area)	7.76	7.11	11.2955
Bias (% total lane area)	-4.54	-0.47	0.7018
C_1	1.31	0.19	0.19
C_2 ($h_{ac} < 5$ in.)	2.1585	0.78	0.78
C_2 ($5 \text{ in.} \leq h_{ac} \leq 12 \text{ in.}$)	$(0.867 + 0.2583 * h_{ac}) * 1$	$(0.867 + 0.2583 * h_{ac}) * 0.26$	$(0.867 + 0.2583 * h_{ac}) * 0.26$

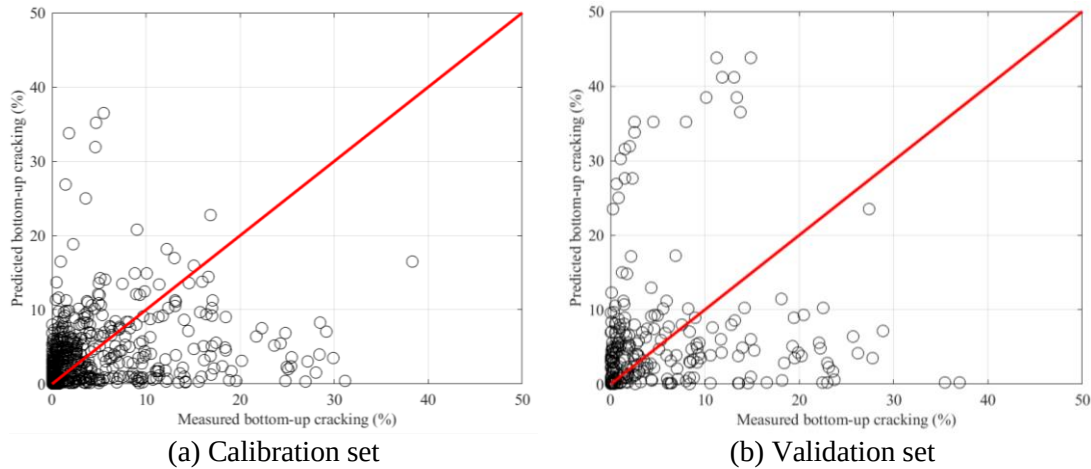


Figure A-3 Local calibration results for bottom-up cracking (split sampling)

Repeated Split Sampling

Like split sampling, repeated split sampling was used with a random split of 70% sections for the calibration set and the remaining 30% for the validation set. This process was repeated 1000 times, where a new random set of calibration and validation sets was picked each time. Repeated split sampling is used to estimate the distribution of different parameters instead of optimizing for a point estimate. Confidence intervals (CI) for each parameter can also be obtained. Tables A-3 to A-5 show the summary for the global model, calibration, and validation sets. It is important to note that coefficient C_2 is a function of total HMA thickness (h_{ac}). For estimating the confidence intervals and distribution of C_2 , it was converted to a single value for all HMA thicknesses. Figures A-4 and A-5 present the distribution of model parameters for calibration and validation sets. In Figures A-4 and A-5, the solid blue line shows the median, the dashed red line shows the mean, the solid black line shows the cumulative distribution and the dashed red lines on both sides show the 2.5th and 97.5th percentiles. The mean SEE is reduced from 8.29 to 7.90 for the calibration and 7.93 for the validation set. Similarly, bias was improved from -4.91 to -0.02 for the calibration and 0.03 for the validation set.

Table A-3 Global model summary (Repeated split sampling)

Parameter	Global model mean	Global model median	Global model lower CI	Global model upper CI
SEE (% total lane area)	8.29	8.29	7.63	8.84
Bias (% total lane area)	-4.91	-4.91	-5.35	-4.47
C_1	1.31	1.31	-	-
C_2 ($h_{ac} < 5$ in.)	2.1585	2.1585	-	-
C_2 (5 in. $\leq h_{ac} \leq 12$ in.)	$(0.867 + 0.2583 * h_{ac}) * 1$	$(0.867 + 0.2583 * h_{ac}) * 1$	-	-

Table A-4 Calibration set summary (Repeated split sampling)

Parameter	Local model mean	Local model median	Local model lower CI	Local model upper CI
SEE (% total lane area)	7.90	7.73	6.49	9.93
Bias (% total lane area)	-0.02	0.00	-0.51	0.42
C ₁	0.26	0.25	0.13	0.42
C ₂ (h _{ac} < 5 in.)	0.60	0.60	0.29	0.89
C ₂ (5 in. ≤ h _{ac} ≤ 12 in.)	(0.867+0.2583* h _{ac})* 0.19	(0.867+0.2583* h _{ac})* 0.19		

Table A-5 Validation set summary (Repeated split sampling)

Parameter	Local model mean	Local model median	Local model lower CI	Local model upper CI
SEE (% total lane area)	7.93	7.68	6.01	10.88
Bias (% total lane area)	0.03	0.02	-2.04	2.27
C ₁	0.26	0.25	0.13	0.42
C ₂ (h _{ac} < 5 in.)	0.60	0.60	0.29	0.89
C ₂ (5 in. ≤ h _{ac} ≤ 12 in.)	(0.867+0.2583* h _{ac})* 0.19	(0.867+0.2583* h _{ac})* 0.19		

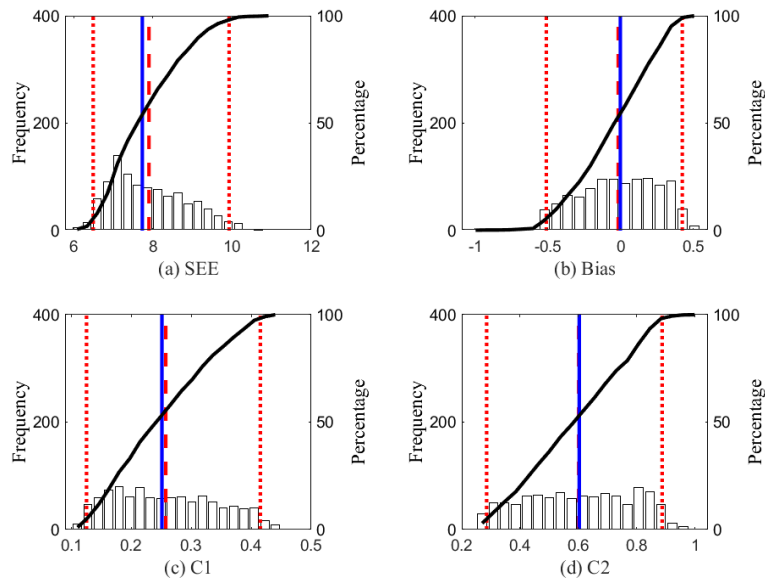


Figure A-4 Local calibration results for bottom-up cracking – calibration dataset (repeated split sampling)

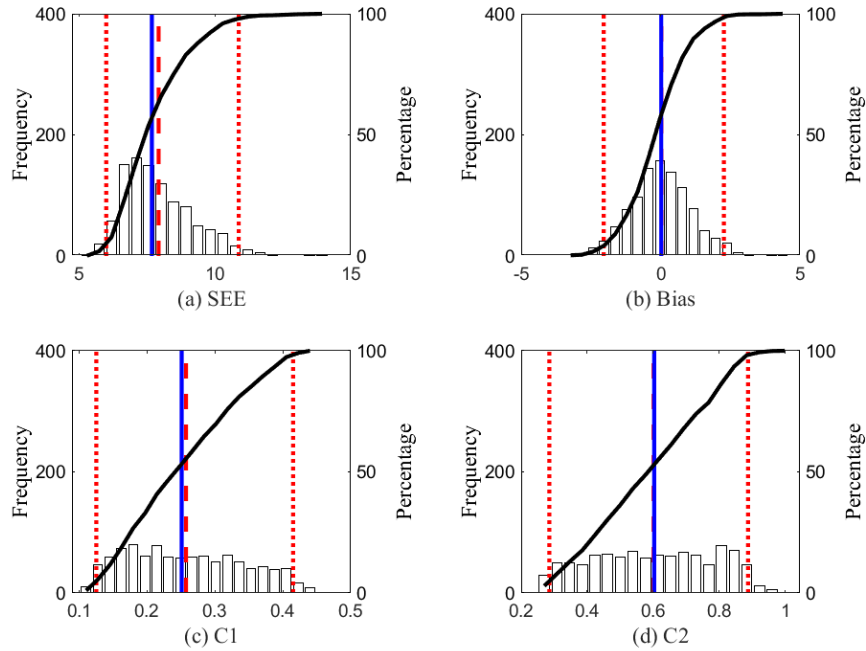


Figure A-5 Local calibration results for bottom-up cracking – validation dataset (repeated split sampling)

Bootstrapping

Bootstrapping was used as a resampling technique to calibrate the bottom-up cracking model. One thousand bootstrap samples were created, randomly sampling with replacement. Unlike repeated split sampling, in bootstrap, the samples were not split; instead, the entire dataset was used. Bootstrapping also generated CI and distribution of model parameters. Tables A-6 and A-7 summarize the model parameters for global and local models, respectively. SEE is slightly increased, whereas bias is significantly improved after local calibration. Figure A-6 shows the distribution of parameters for the 1000 bootstrap samples.

Table A-6 Bootstrapping global model summary

Parameter	Global model mean	Global model median	Global model lower CI	Global model upper CI
SEE (% total lane area)	8.30	8.30	7.38	9.20
Bias (% total lane area)	-4.91	-4.91	-5.53	-4.33
C_1	1.31	1.31	-	-
C_2 ($h_{ac} < 5$ in.)	2.1585	2.1585	-	-
C_2 (5 in. $\leq h_{ac} \leq 12$ in.)	$(0.867 + 0.2583 * h_{ac}) * 1$	$(0.867 + 0.2583 * h_{ac}) * 1$	-	-

Table A-7 Bootstrapping local calibration results summary

Parameter	Local model mean	Local model median	Local model lower CI	Local model upper CI
SEE (% total lane area)	8.73	8.30	6.21	12.83
Bias (% total lane area)	0.00	-0.03	-0.80	0.68
C_1	0.23	0.20	0.01	0.54
C_2 ($h_{ac} < 5$ in.)	0.70	0.73	0.04	1.29
C_2 (5 in. $\leq h_{ac} \leq 12$ in.)	$(0.867+0.2583 \cdot h_{ac})^* 0.22$	$(0.867+0.2583 \cdot h_{ac})^* 0.23$		

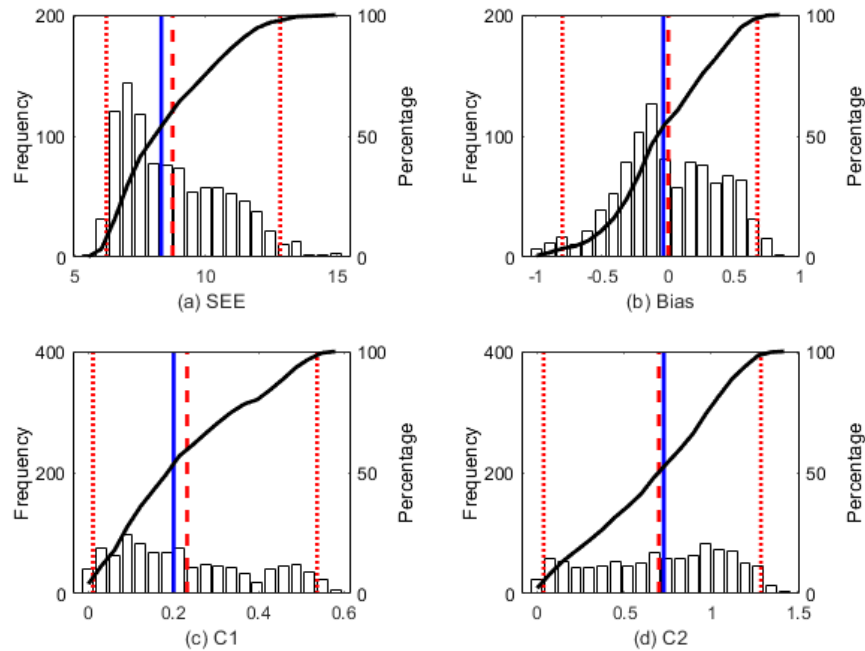


Figure A-6 Local calibration results for bottom-up cracking (bootstrapping)

Summary

All calibration approaches have significantly improved the bottom-up cracking model. Table A-8 shows the summary of all sampling techniques. It should be noted that these calibrations were performed with specific limits on the calibration coefficients taken from the literature, as mentioned in Chapter 2. These limits ensure that we get reasonable and practical calibration results.

Table A-8 Summary of results for all sampling techniques (Option a)

Sampling technique	SEE	Bias	C_1	C_2 ($h_{ac} < 5$ in.)	C_2 (5 in. $\leq h_{ac} \leq 12$ in.)
No sampling	8.08	0.17	0.22	0.66	$(0.867+0.2583 \cdot h_{ac})^* 0.21$
Split sampling	7.11	-0.47	0.19	0.78	$(0.867+0.2583 \cdot h_{ac})^* 0.26$
Repeated split sampling	7.90	-0.02	0.26	0.60	$(0.867+0.2583 \cdot h_{ac})^* 0.20$
Bootstrapping	8.73	0.00	0.23	0.70	$(0.867+0.2583 \cdot h_{ac})^* 0.22$

TOP-DOWN CRACKING MODEL

The following section shows the calibration of the top-down cracking model. The model contains crack initiation and crack propagation models. Since the actual crack initiation time is not known, it was not possible to calibrate the crack initiation model separately. So, a single function was used by substituting the crack initiation function with the crack propagation function. Initially, an attempt was made to change all eight coefficients simultaneously. This approach had some challenges:

- The model has some mathematical limitations. High values for C_3 give mathematical errors in the Pavement-ME output.
- There is no current literature available for the top-down cracking model. Therefore, estimating the range for each coefficient to be used in optimization was difficult.
- The model has numerous coefficients with coefficient values ranging from 0.011 to 64271618. This makes the optimization challenging to converge.

The top-down cracking model was calibrated in Microsoft Excel by combining engineering judgment and the solver function. Four coefficients from the crack initiation function ($kL2$, $kL3$, $kL4$, $kL5$) and two from the crack propagation function (C_1 , C_2) have been calibrated. No sampling method was used for this calibration.

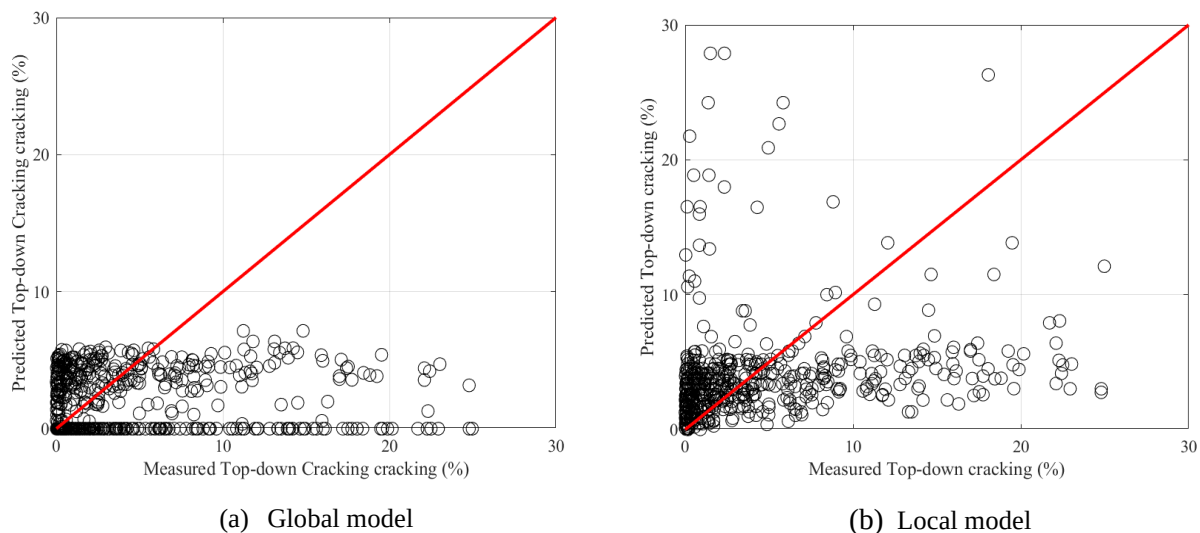


Figure A-7 Predicted vs. measured top-down cracking (No sampling)

Figure A-7 shows the predicted vs. measured top-down cracking, and Figure A-8 shows the predicted and measured top-down cracking with time. The predicted and measured top-down cracking does not follow similar trends. Most top-down cracking predictions are limited to a specific time series curve. Table A-9 summarizes model parameters. The SEE and bias are improved. The reliability of the top-down cracking model is estimated by developing a relationship between the standard deviation of the measured cracking, and the mean predicted cracking. Table A-10 outlines the standard error equations for the global and calibrated model.

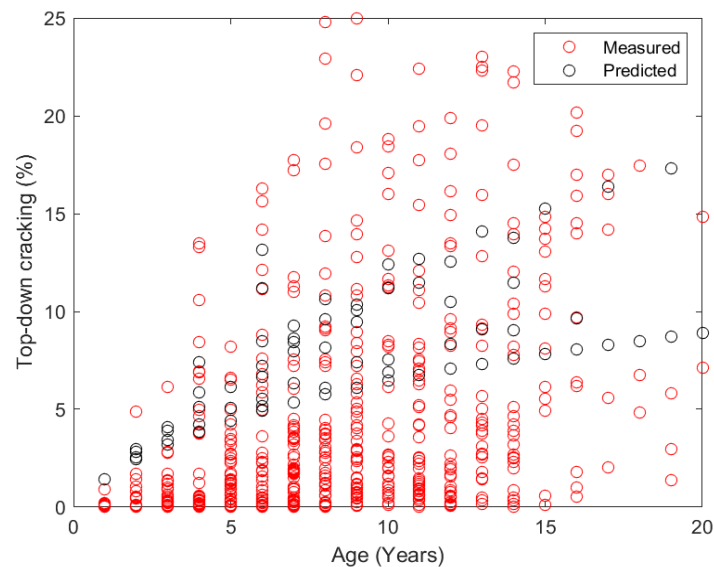


Figure A-8 Measured and predicted top-down-cracking (time series)

Table A-9 Calibration results for top-down cracking

Parameters	Global model	Local model
SEE	6.37	5.59
Bias	-2.36	1.60
K_{L2}	0.2855	0.90
K_{L3}	0.011	0.09
K_{L4}	0.01488	0.101
K_{L5}	3.266	3.260
C_1	2.5219	0.30
C_2	0.8069	1.155

Table A-10 Reliability equation for top-down cracking

Pavement-ME model	Global model equation	Local model equation
Top-down cracking	$s_{e(Top-down)} = 0.3657 \times TOP + 3.6563$	$s_{e(Top-down)} = 0.6417 \times TOP + 0.5014$

THERMAL CRACKING MODEL

The thermal cracking model was calibrated for Level 1 inputs in the Pavement-ME. The model calibration only considered sections with Performance Grade (PG) binder type. The thermal cracking model was calibrated as a single K -value by running Pavement-ME multiple times. Although calibration coefficient K is a function of mean annual air temperature (MAAT), it was calibrated as a single value similar to the previous version of Pavement-ME (version 2.3). For this purpose, the Pavement-ME was run at different K values (0.25, 0.65, 0.75, 0.85, 0.95 and 1.35). SEE and bias were determined for each value of K . Table A-11 summarizes the SEE and bias for the global model and different K values. Based on the SEE and bias, a value of 0.85 is recommended. Recalibration improved the SEE and bias, but thermal cracking predictions still show high variability. Figure A-9 shows the predicted vs. measured thermal cracking for the global and local models at $K=0.85$. As previously explained in Chapter 3, measured thermal cracking values have been capped at 2112 feet/mile. This means any measured value of more than 2112 feet/mile for sections has been removed from the calibration data.

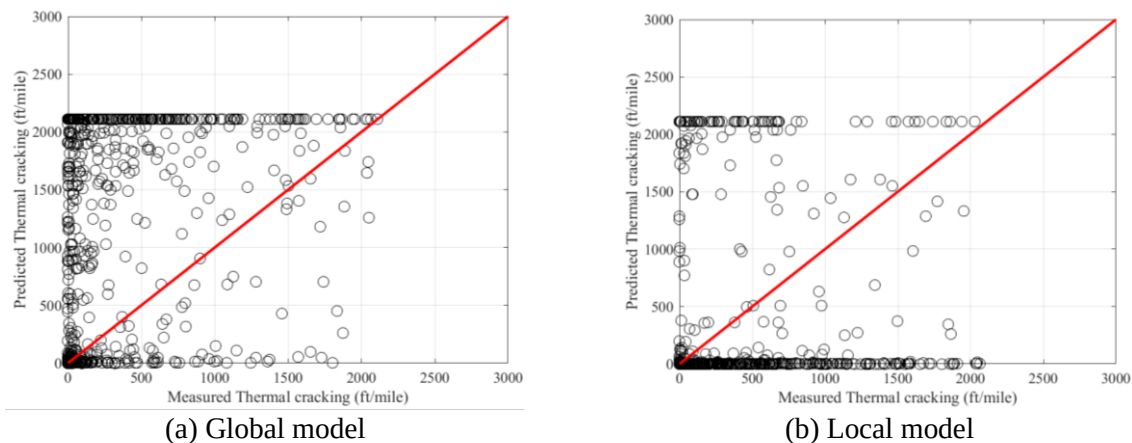


Figure A-9 Predicted vs. measured thermal cracking (at $K=0.85$)

Table A-11 Thermal cracking calibration results

Parameter	SEE	Bias
Global model	1225	-812
$K = 0.25$	650	272
$K = 0.65$	760	172
$K = 0.75$	813	106
$K = 0.85$	851	20
$K = 0.95$	893	-71
$K = 1.35$	1077	-471

The standard error equations were developed using the standard deviation of the measured cracking and mean predicted cracking, as explained in Chapter 4. Table A-12 summarizes the standard error equations for the global and locally calibrated models.

Table A-12 Reliability summary for thermal cracking

Pavement-ME model	Global model equation	Local model equation
Thermal cracking	$s_e = 0.14(TC) + 168$	$s_e = 0.1223(TC) + 400.9$

RUTTING MODEL (METHOD 1)

No Sampling

Pavement-ME predictions for individual layer rutting were matched against measured rutting determined by using the transverse profile analysis results, as discussed in Chapter 4. Figures A-10 to A-12 show the predicted vs. measured rutting for AC, base, and subgrade layers, respectively. The Pavement-ME under-predicts AC rutting and over-predicts base and subgrade rutting. Table A-13 shows the SEE and bias, whereas Table A-14 shows the calibrated coefficients. Both SEE and bias significantly improved for all layers.

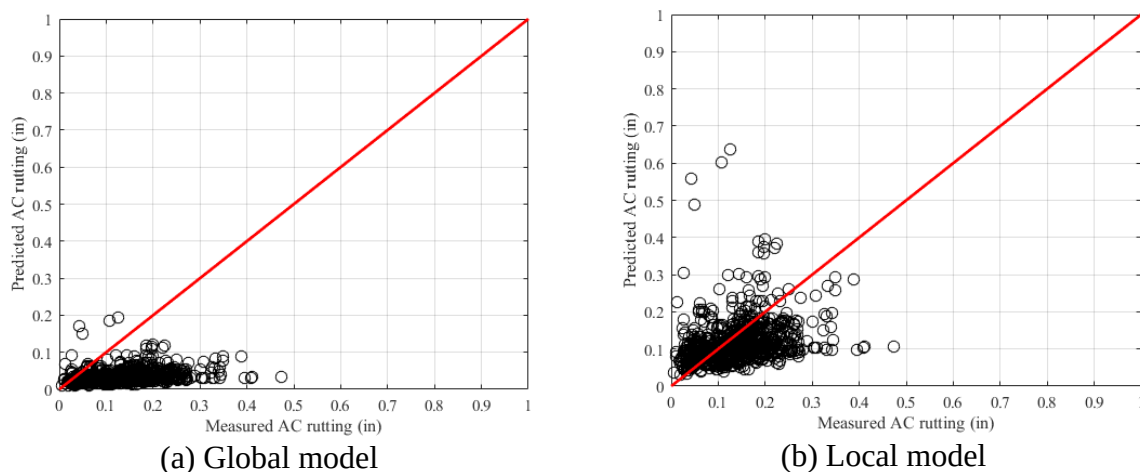
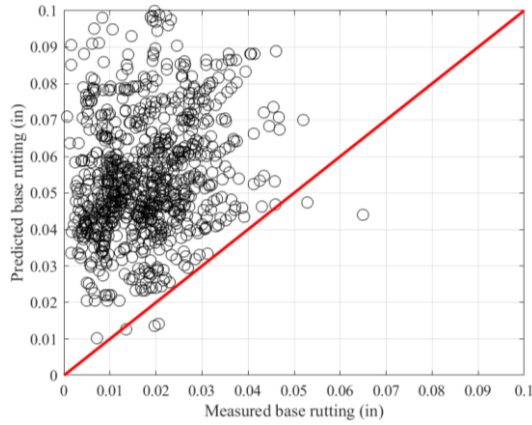
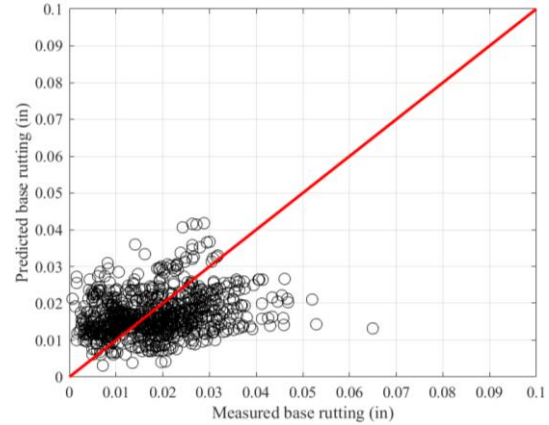


Figure A-10 Predicted vs. measured AC rutting (No sampling)

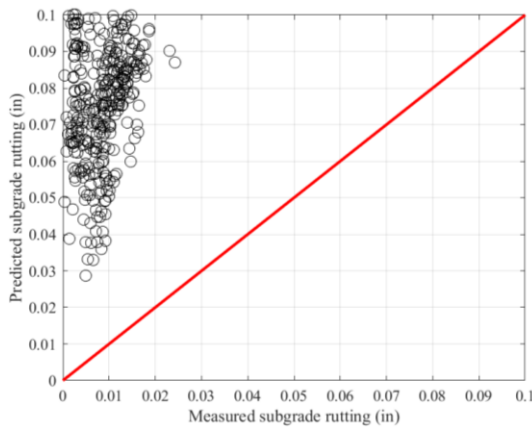


(a) Global model

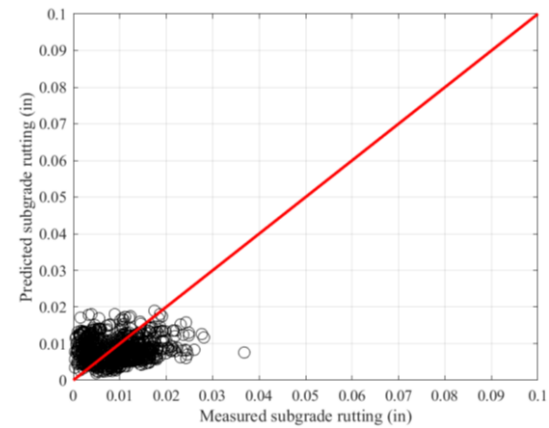


(b) Local model

Figure A-11 Predicted vs. measured base rutting (No sampling)



(a) Global model



(b) Local model

Figure A-12 Predicted vs. measured subgrade rutting (No sampling)

Table A-13 Rutting models SEE and bias

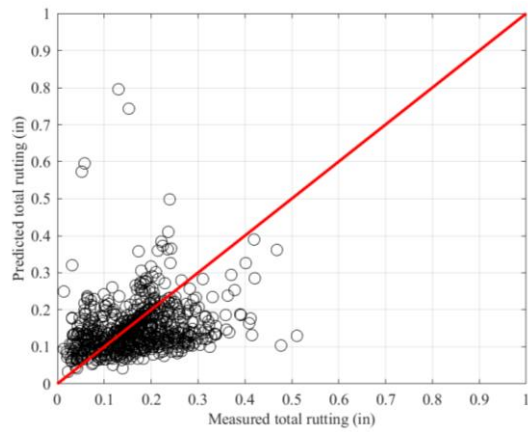
Layer	Global model		Local model	
	SEE (in.)	Bias (in.)	SEE (in.)	Bias (in.)
HMA rut	0.2579	0.2015	0.0812	-0.0138
Base rut	0.0426	0.0380	0.0099	-0.0011
Subgrade	0.1184	0.1095	0.0062	-0.0009

Table A-14 Rutting model calibration coefficients

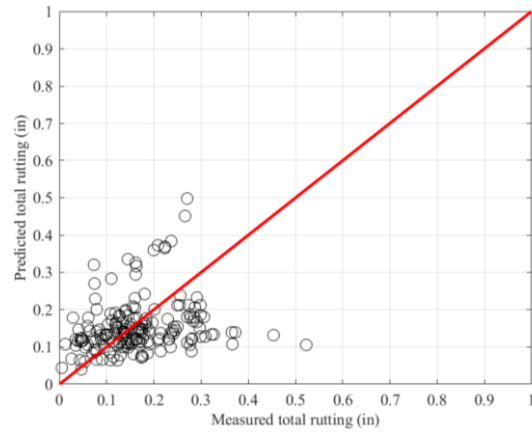
Calibration coefficient	Global model	Local model
HMA rutting (br1)	0.4	0.1466
Base rutting (bs1)	1.0000	0.3003
Subgrade rutting (bsg1)	1.0000	0.0691

Split Sampling

Split sampling was performed on 70% of the sections for the calibration set and 30% for the validation set. Figures A-13 to A-15 show the predicted vs. measured for calibration and validation set for different layers. All layers show reasonable validation results.

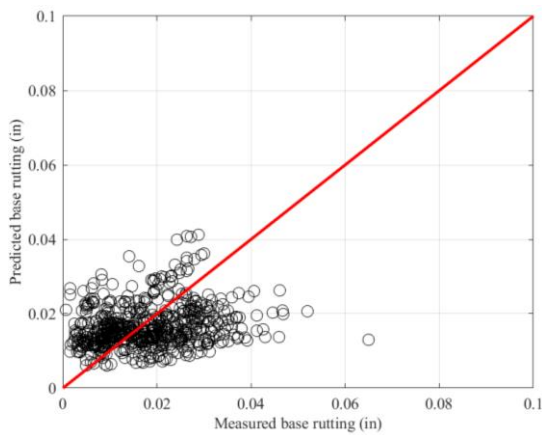


(a) Calibration set

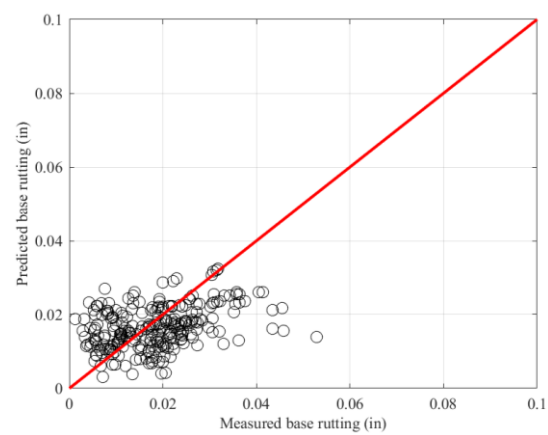


(b) Validation set

Figure A-13 Predicted vs. measured AC rutting (Split sampling)



(a) Calibration set



(b) Validation set

Figure A-14 Predicted vs. measured Base rutting (Split sampling)

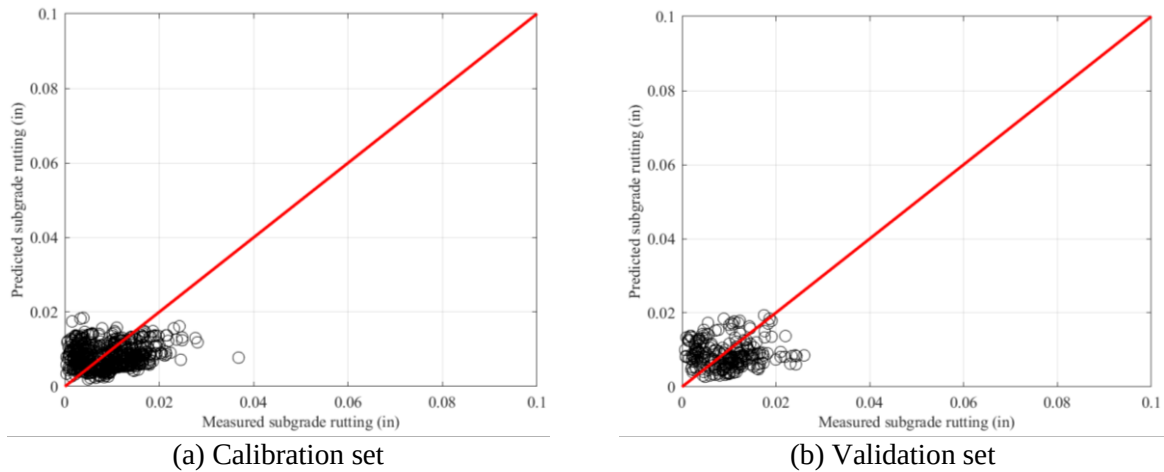


Figure A-15 Predicted vs. measured Subgrade rutting (Split sampling)

Table A-15 shows the SEE, bias, and model parameters for the global model, and Table A-16 shows the same for the calibration-validation set. Both SEE and bias significantly improved for all layers.

Table A-15 Rutting global model results

Layer	SEE	Bias	Coefficient
HMA rut	0.2454	0.1759	0.4
Base rut	0.0872	-0.0138	1.0000
Subgrade	0.1153	0.1071	1.0000

Table A-16 Rutting local model results

Layer	Calibration set			Validation set		
	SEE	Bias	Coefficient	SEE	Bias	Coefficient
HMA rut	0.0962	-0.0165	0.0705	0.1008	-0.0117	0.0705
Base rut	0.0102	-0.0012	0.2955	0.0092	-0.0018	0.2955
Subgrade	0.0061	-0.0008	0.0705	0.0064	-0.0007	0.0705

Repeated Split Sampling

Repeated split sampling was performed for 1000 split samples with new calibration and validation sets. Figures A-16 to A-18 show the distribution of model parameters for calibration and validation set for different layers. Tables A-17 to A-19 show the SEE, bias, model parameters, CI for the global model, and the calibration and validation sets, respectively. The rutting model significantly improved after local calibration.

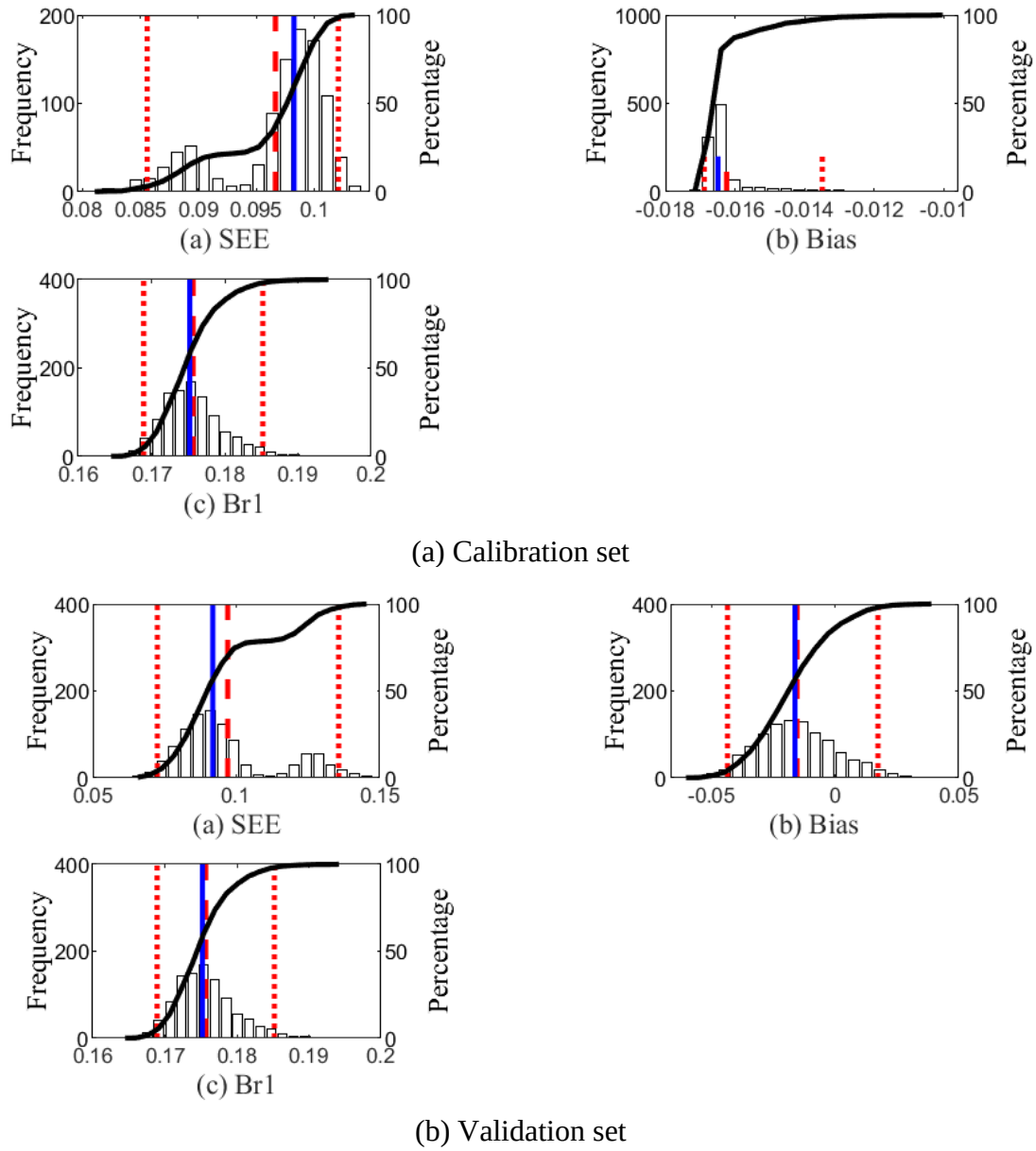


Figure A-16 Distribution of calibration parameters - AC rutting (Repeated split sampling)

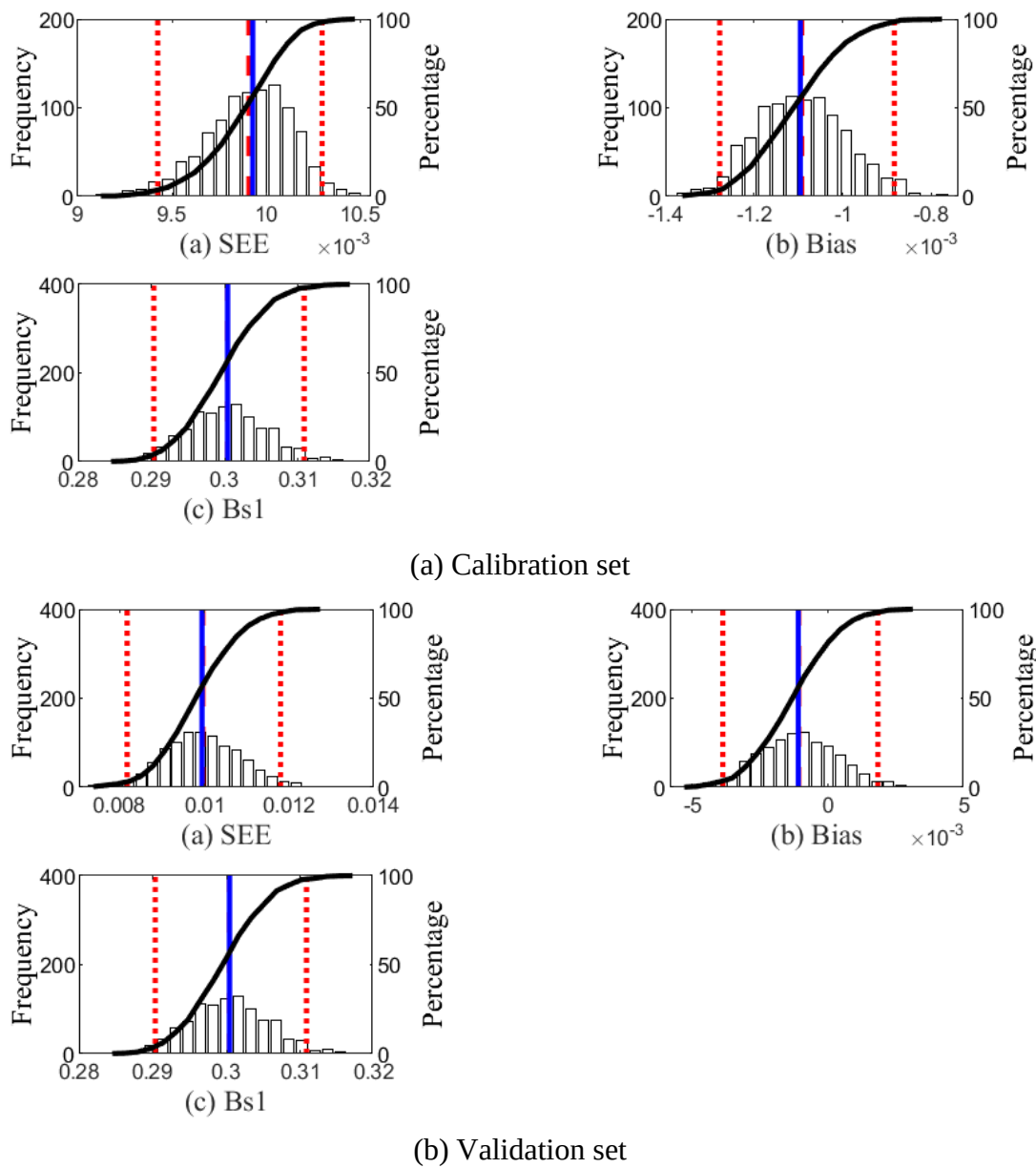


Figure A-17 Distribution of calibration parameters - Base rutting (Repeated split sampling)

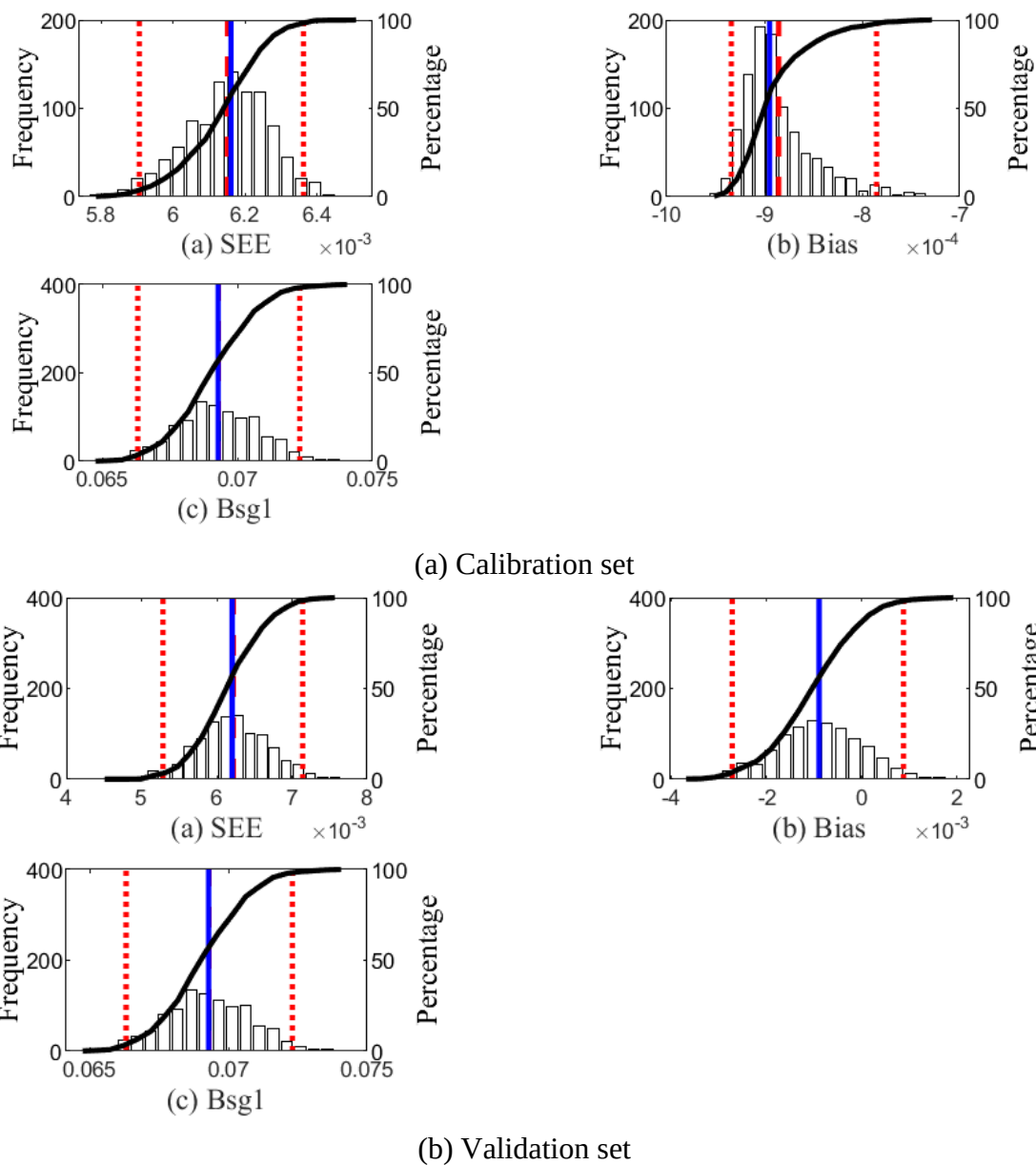


Figure A-18 Distribution of calibration parameters - Subgrade rutting (Repeated split sampling)

Table A-17 Global model results (repeated split sampling)

Layer	Average SEE	SEE Lower CI	SEE Upper CI	Average bias (in.)	Bias Lower CI	Bias Upper CI
HMA	0.2387	0.2097	0.2540	0.1743	0.1617	0.1853
Base	0.0426	0.0409	0.0440	0.0380	0.0367	0.0394
Subgrade	0.1185	0.1150	0.1216	0.1095	0.1064	0.1126

Table A-18 Local model calibration results (repeated split sampling)

Statistics	HMA rutting	Base rutting	Subgrade rutting
Average SEE	0.0966	0.0099	0.0062
SEE Lower CI	0.0856	0.0094	0.0059
SEE Upper CI	0.1021	0.0103	0.0064
Average bias (in.)	-0.0162	-0.0011	-0.0009
Bias Lower CI	-0.0169	-0.0013	-0.0009
Bias Upper CI	-0.0135	-0.0009	-0.0008
Average calibration coefficient	0.1757	0.3003	0.0693
Calibration coefficient Lower CI	0.1689	0.2897	0.0663
Calibration coefficient Upper CI	0.1852	0.3115	0.0723

Table A-19 Local model validation results (repeated split sampling)

Statistics	HMA rutting	Base rutting	Subgrade rutting
Average SEE	0.0971	0.0100	0.0062
SEE Lower CI	0.0725	0.0084	0.0053
SEE Upper CI	0.1358	0.0119	0.0071
Average bias (in.)	-0.0153	-0.0011	-0.0009
Bias Lower CI	-0.0434	-0.0041	-0.0027
Bias Upper CI	0.0174	0.0017	0.0009
Average calibration coefficient	0.1757	0.3003	0.0693
Calibration coefficient Lower CI	0.1689	0.2897	0.0663
Calibration coefficient Upper CI	0.1852	0.3115	0.0723

Bootstrapping

Bootstrapping was performed with 1000 bootstrap samples with replacement. Figures A-19 to A-21 show the distribution of model parameters for AC, base, and subgrade rutting. Tables A-20 and A-21 summarize the calibration results for the global and local models. Model parameter distribution and CI provide a more reliable estimate of model coefficients. Moreover, SEE and bias significantly improved for all layers.

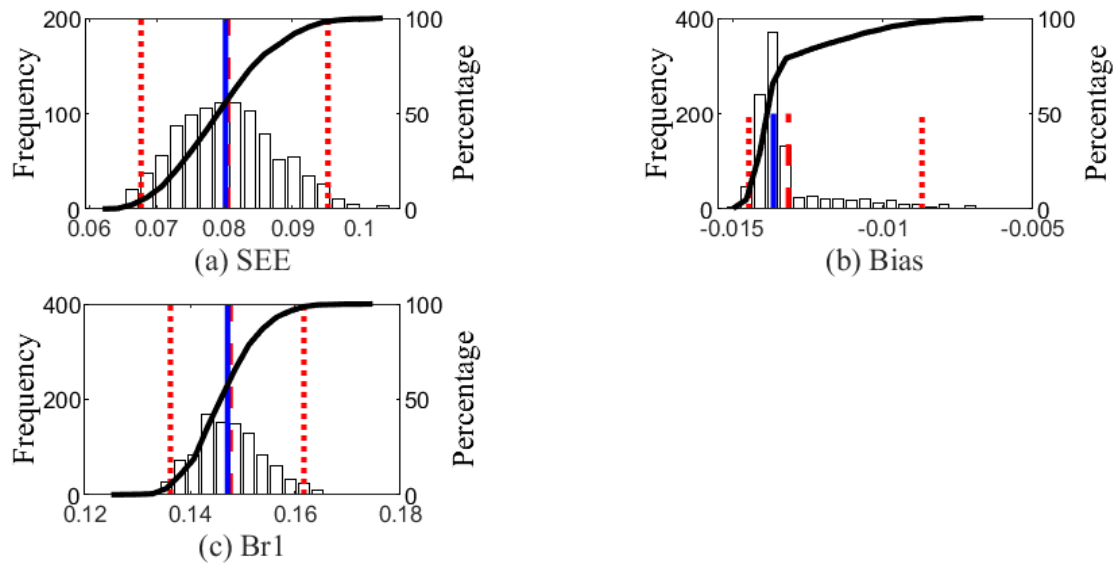


Figure A-19 Distribution of calibration parameters - AC rutting (Bootstrapping)

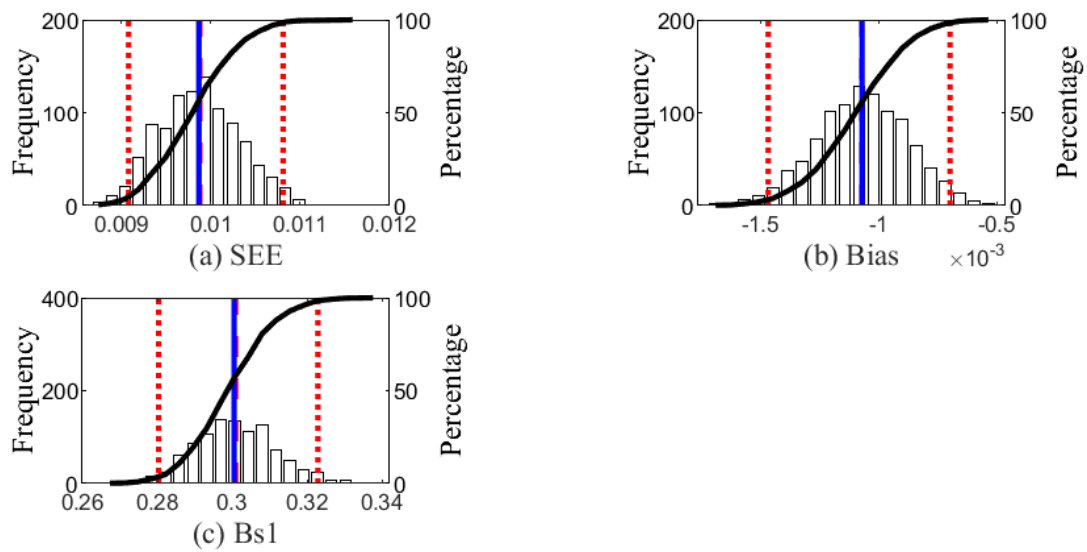


Figure A-20 Distribution of calibration parameters - Base rutting (Bootstrapping)

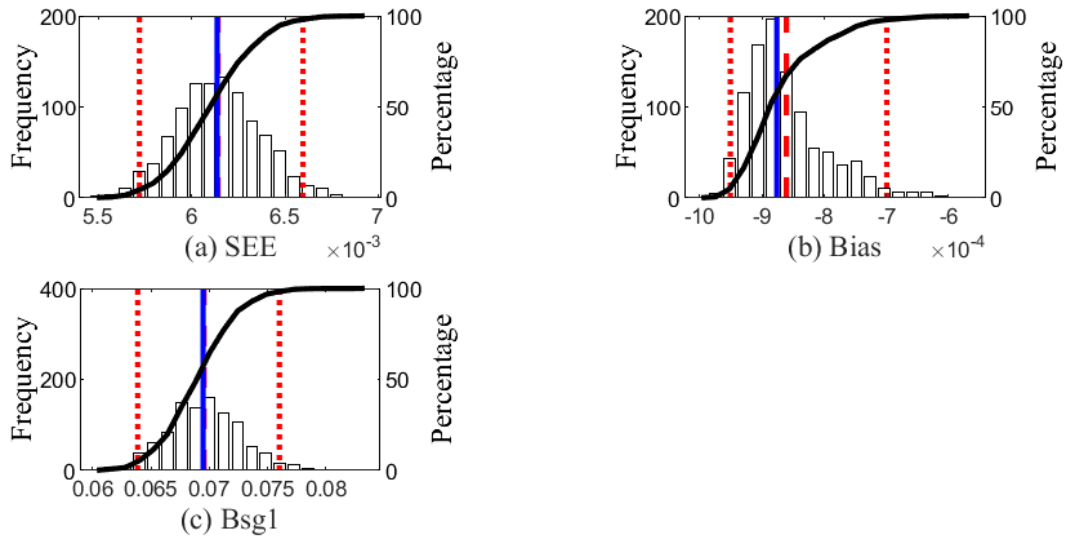


Figure A-21 Distribution of calibration parameters-Subgrade rutting (Bootstrapping)

Table A-20 Global rutting model summary (Bootstrapping)

Layer type	Average SEE	SEE Lower CI	SEE Upper CI	Average bias (in.)	Bias Lower CI	Bias Upper CI
HMA	0.2565	0.2174	0.3047	0.2010	0.1796	0.2238
Base	0.0425	0.0396	0.0456	0.0380	0.0355	0.0408
Subgrade	0.1183	0.1117	0.1251	0.1094	0.1032	0.1159

Table A-21 Local rutting model summary (Bootstrapping)

Statistics	HMA rutting	Base rutting	Subgrade rutting
Average SEE	0.0805	0.0099	0.0061
SEE Lower CI	0.0677	0.0091	0.0057
SEE Upper CI	0.0953	0.0108	0.0066
Average bias (in.)	-0.0131	-0.0011	-0.0009
Bias Lower CI	-0.0145	-0.0015	-0.0010
Bias Upper CI	-0.0087	-0.0007	-0.0007
Average calibration coefficient	0.1476	0.3009	0.0696
Calibration coefficient Lower CI	0.1363	0.2803	0.0639
Calibration coefficient Upper CI	0.1616	0.3228	0.0760

Summary

Results for Method 1 are summarized in Table A-22. All calibration approaches have improved the SEE and bias. Bootstrap shows the lowest SEE and bias for all layers.

Table A-22 Rutting model calibration results summary

Sampling Technique	Pavement layer rutting	SEE	Bias	Calibration coefficient
No sampling	HMA	0.0812	-0.0138	0.1466
	Base	0.0099	-0.0011	0.3003
	Subgrade	0.0062	-0.0009	0.0691
Split sampling	HMA	0.0962	-0.0165	0.0705
	Base	0.0102	-0.0012	0.2955
	Subgrade	0.0061	-0.0008	0.0705
Repeated split sampling	HMA	0.0971	-0.0153	0.1757
	Base	0.0099	-0.0011	0.3003
	Subgrade	0.0062	-0.0009	0.0693
Bootstrapping	HMA	0.080	-0.013	0.148
	Base	0.010	-0.001	0.301
	Subgrade	0.006	-0.001	0.070

JOINT FAULTING MODEL

The calibration of the faulting model was performed using the CAT tool. No sampling technique was used for the calibration. In the first step, the most sensitive coefficients, C_1 and C_6 , were simultaneously calibrated. In the next step, C_1 and C_6 were kept at the calibrated value, and C_2 was calibrated. All other coefficients (C_3 , C_4 , C_5 , C_7 , and C_8) were kept at the global values. It should be noted that the measured faulting was cut to 0.4 inches, as mentioned in Chapter 3. Figure A-22 shows the predicted vs. measured joint faulting for the global and local models. Figure A-23 shows the measured and predicted joint faulting with time. In Figure A-23, the predicted faulting is in the same range as measured faulting except for high values for measured faulting.

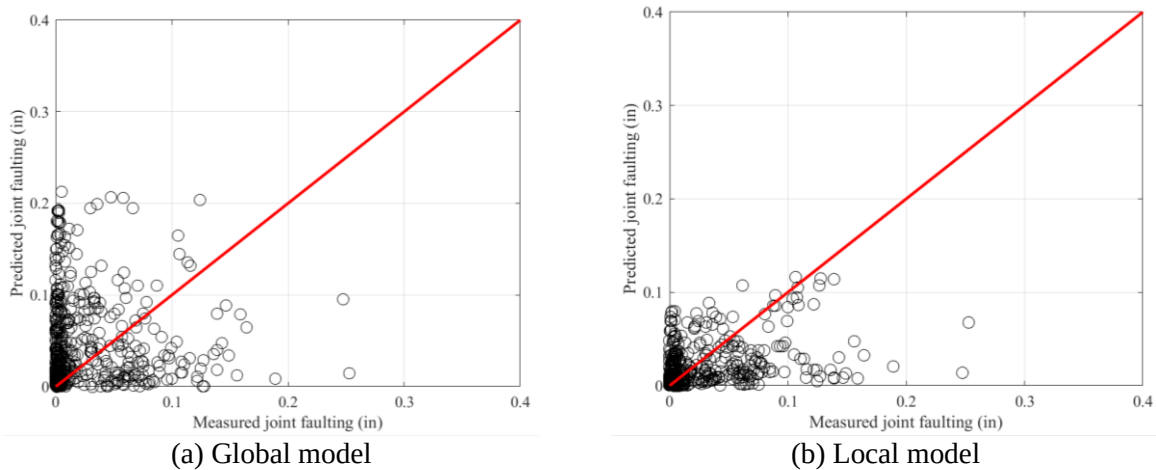


Figure A-22 Calibration results for joint faulting

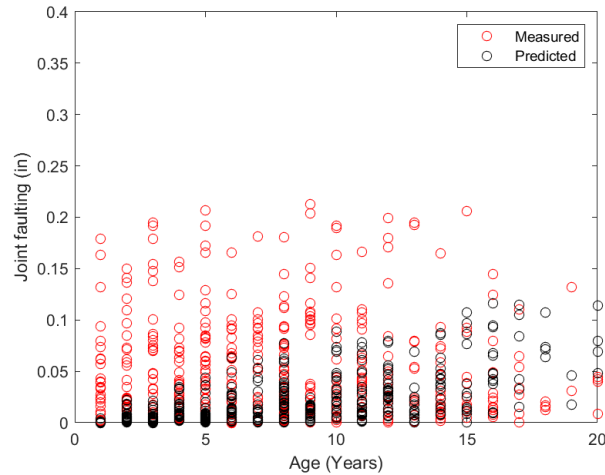


Figure A-23 Measured and predicted joint faulting (time series)

Table A-23 summarizes local calibration and the corresponding model parameters. SEE and bias are significantly improved.

Table A-23 Summary of faulting model calibration

Parameter	Global model	Local model
SEE	0.06	0.03
Bias	0.01	0.00
C ₁	0.595	0.8
C ₂	1.636	1.3889
C ₃	0.00217	0.00217
C ₄	0.00444	0.00444
C ₅	250	250
C ₆	0.47	0.2
C ₇	7.3	7.3
C ₈	400	400

The standard error equations were estimated, establishing a relationship between the standard deviation of the measured faulting and mean predicted faulting, as explained in Chapter 4. Table A-24 summarizes standard error equations for the faulting model.

Table A-24 Faulting model reliability

Pavement-ME model	Global model equation	Local model equation
Joint faulting	$s_{e(Fault)} = 0.07162(Fault)^{0.368} + 0.00806$	$s_{e(Fault)} = 0.0902(Fault)^{0.2038}$