

POLITECHNIKA WARSZAWSKA

DYSCYPLINA NAUKOWA INŻYNIERIA ŚRODOWISKA, GÓRNICTWO
I ENERGETYKA /
DZIEDZINA NAUK INŻYNIERYJNO-TECHNICZNYCH

Rozprawa doktorska

mgr inż. Paweł Domitr

**Metodyka wykonywania obliczeń najlepszego szacowania wraz
z oceną niepewności na podstawie metod wstecznej kwantyfikacji
niepewności**

Promotor
dr hab. inż. Rafał Laskowski

WARSZAWA 2023

Podziękowania

Na wstępie chciałbym podziękować mojemu promotorowi doktorowi hab. inż. Rafałowi Laskowskiemu za pomoc, wsparcie i stworzenie pozytywnej atmosfery współpracy.

Równie wielkie podziękowania składam mojemu opiekunowi z ramienia Państwowej Agencji Atomistyki doktorowi inż. Ernestowi Staroniowi, który od ponad dziesięciu lat jest moim mentorem i nieocenionym wsparciem w pracy zawodowej.

Składam również serdeczne podziękowania Profesorowi Tomaszowi Wiśniewskiemu, który pełnił funkcję mojego opiekuna naukowego przez rok, od początku motywując mnie i zachęcając do osiągania wymiernych postępów w pracach nad rozprawą doktorską.

Dziękuję moim kolegom i koleżankom z Państwowej Agencji Atomistyki, w szczególności Mateuszowi Włostowskiemu za wartościowe dyskusje oraz wszelką pomoc.

Dziękuję również za znakomitą współpracę moim kolegom z Instytutu Techniki Ciepłej doktorowi Piotrowi Darnowskiemu oraz doktorowi Michałowi Spirzewskiemu.

Specjalne podziękowania składam moim Rodzicom, bez których wsparcia nie miałbym możliwości napisania tej pracy.

Z całego serca dziękuję mojej Żonie Małgorzacie za niekończące się wsparcie, motywację, cierpliwość oraz wiarę we mnie.

Prace zrealizowano w ramach programu „Doktorat wdrożeniowy” przy współpracy z Państwową Agencją Atomistyki.

Streszczenie rozprawy doktorskiej w języku polskim

Niniejsza rozprawa traktuje o nowej metodyce wykonywania obliczeń najlepszego szacowania wraz z oceną niepewności, która pozwala na usprawnienie dotychczas stosowanych metod poprzez uwzględnienie wstecznej kwantyfikacji niepewności parametrów wejściowych. Metodyka najlepszego szacowania wraz z oceną niepewności może być wykorzystywana jako jedno z podejść do wykonywania deterministycznych analiz bezpieczeństwa dla awarii projektowych w reaktorach jądrowych. Zaproponowana metodyka bazuje na połączeniu najlepszych praktyk wykonywania obliczeń niepewności przy wykorzystaniu dotychczasowo stosowanych metod, jednocześnie pozwalając na minimalizację niepewności wynikających z efektu użytkownika. Równocześnie zaproponowano, wdrożono oraz zademonstrowano w praktyce użyteczność dwóch metod obliczeniowych wykonywania wstecznej kwantyfikacji niepewności. Uzyskane w wyniku przeprowadzonej kwantyfikacji rozkłady gęstości parametrów zostały następnie poddane procesowi walidacji.

W rozprawie przedstawiono najważniejsze zagadnienia związane z wykonywaniem deterministycznych analiz bezpieczeństwa stosując podejście najlepszego szacowania wraz z oceną niepewności. Omówiono obecnie stosowane metody wraz z ich zaletami i wadami, co stanowiło podstawę i motywację do zaproponowania usprawnień i opracowania nowej metodyki. W dalszej części rozprawy przedstawiono opis metodyki składającej się z sześciu głównych elementów. Omówiono szczegółowo wszystkie elementy wraz ze składającymi się na nie krokami. Kolejne części rozprawy przedstawiają wdrożenie każdego elementu metodyki. Znacząca część rozprawy poświęcona jest opisowi metod obliczeniowych wstecznej kwantyfikacji niepewności parametrów wejściowych kodu cieplno-przepływowego TRACE opisujących trzy wybrane zjawiska obserwowane podczas awarii typu LBLOCA: wpływ krytyczny, wymianę ciepła w rdzeniu oraz propagację frontu zalewania. Obliczenia niepewności dla tych parametrów wykonano dla dwóch instalacji eksperymentalnych Marviken i Flecht Seaset. Pierwsza z metod oparta jest o wnioskowanie Bayesowskie oraz wykorzystanie numerycznych metod Monte Carlo. Druga z metod oparta jest o algorytm uczenia maszyn. Algorytm nadzorowanego uczenia maszyn zaimplementowano do oceny miary rozbieżności między danymi eksperymentalnymi a obliczeniami wykonywanymi z zastosowaniem kodu obliczeniowego. Rozwiązaniem problemu wstecznej kwantyfikacji niepewności parametrów wejściowych kodu obliczeniowego są rozkłady gęstości prawdopodobieństwa tych parametrów oparte o wiedzę na temat ich wpływu na uzyskiwane wyniki obliczeń w odniesieniu do danych eksperymentalnych. W rozprawie przedstawiono następnie obliczenia walidacyjne dla

uzyskanych rozkładów. Walidację wykonano przeprowadzając obliczenia dla instalacji eksperymentalnej LOFT, w której badano zjawiska fizyczne zachodzące podczas awarii typu LBLOCA. Zaprezentowano również ostateczne obliczenia najlepszego szacowania wraz z oceną niepewności dla ciśnieniowego reaktora wodnego dla awarii typu LBLOCA. Obliczenia te stanowią ostatni element zaproponowanej metodyki. Rozprawa zakończona jest przedstawieniem wniosków na temat użyteczności metodyki, możliwości ograniczania efektu użytkownika, jakie dają zaproponowane metody wstecznej kwantyfikacji niepewności oraz dalszego kierunku badań nad rozwojem metodyki i metod obliczeniowych.

Słowa kluczowe: energetyka jądrowa, bezpieczeństwo jądrowe, obliczenia niepewności, wsteczna kwantyfikacja niepewności, TRACE, Markov Chain Monte Carlo, uczenie maszyn

Abstract

The doctoral dissertation is devoted to the development of a new methodology of best estimate plus uncertainty calculations, which will be an improvement over the currently implemented methodologies by introducing inverse uncertainty quantification of input parameters. Best estimate plus uncertainty methodologies may be used as one of the approaches to perform deterministic safety analyses of design basis accidents in nuclear reactors. The proposed methodology is based on best practices obtained from the implementation of current methods, additionally allowing for a reduction of the user effect. Development of the methodology leads to the proposition, implementation and demonstration of two computational methods of inverse uncertainty quantification. Probability density functions of input parameters which are the results of inverse quantification were validated in the next steps of the developed methodology.

The doctoral dissertation opens with a discussion of the most important issues of deterministic safety analyses performed with best estimate plus uncertainty approach. Two currently used methods have been discussed along with a presentation of their strengths and weaknesses, which formed a basis and motivation for improvements proposition and development of a new methodology. The following chapters of the dissertation describe the methodology, consisting of six elements. All six elements were thoroughly discussed. Implementation of each element of the methodology followed. The thesis's central part consists of a description of two inverse uncertainty quantification methods. The methods were used to quantify the uncertainties of input parameters of the thermal-hydraulic code TRACE. Those input parameters characterize three phenomena observed during LBLOCA: critical flow, heat transfer in the core and quench front propagation. Inverse quantification for those parameters was performed for two separate effect tests: Marviken critical flow experiment and Flecht Seaset reflood experiment. First of the developed methods is based on Bayesian inference and the application of Monte Carlo methods for random sampling. The second method is based on machine learning algorithms. A supervised machine learning algorithm based on random forests was implemented to quantify the discrepancy between experimental data and the results of calculations performed with the TRACE computer code. The solution to the inverse uncertainty quantification problem are probability density functions of code input parameters, inferred by their influence on the output parameters of calculations and their agreement with experimental data. Probability density functions were validated in the next step of the methodology presented in the dissertation. Validation was performed for the LOFT facility, where an array of phenomena occurring during LBLOCA were studied. Finally, best estimate plus uncertainty calculations using proposed

probability density functions for input parameters were performed for a Pressurized Water Reactor for a LBLOCA scenario. The dissertation ends with a closing discussion, conclusions on the effectiveness and utility of the proposed methodology, and user effect reduction achievement by the introduction of inverse uncertainty quantification methods. Further research on both the methodology and computational methods is also proposed.

Keywords: nuclear power, nuclear safety, best estimate plus uncertainty, inverse uncertainty quantification, TRACE, Markov Chain Monte Carlo, machine learning

Spis treści

1.	Wprowadzenie	11
1.1.	Deterministyczne analizy bezpieczeństwa dla reaktorów jądrowych	11
1.2.	Zagadnienie niepewności w deterministycznych analizach bezpieczeństwa	15
1.3.	Motywacja rozprawy	18
1.4.	Cel i zakres rozprawy	19
2.	Metody obliczeń najlepszego szacowania wraz z oceną niepewności - stan wiedzy	21
2.1.	Obecnie stosowane metody najlepszego szacowania wraz z oceną niepewności w deterministycznych analizach bezpieczeństwa.....	22
2.1.1.	Metoda propagacji niepewności wejściowych	22
2.1.2.	Metoda propagacji niepewności wyjściowych	24
2.1.3.	Metody wykorzystujące analizę Bayesowską oraz globalne analizy wrażliwości	25
2.2.	Ograniczenia obecnie stosowanych metod.....	27
2.2.1.	Benchmark BEMUSE	28
2.2.2.	Benchmark PREMIUM	32
3.	Opis problemu badawczego – metodyka wykonywania obliczeń najlepszego szacowania wraz z oceną niepewności.....	36
3.1.	Element pierwszy metodyki – specyfikacja problemu kwantyfikacji niepewności wejściowych.....	38
3.2.	Element drugi metodyki – przygotowanie odpowiedniej bazy eksperymentalnej	40
3.3.	Element trzeci metodyki – opracowanie i ocena modelu instalacji eksperymentalnej.....	41
3.4.	Element czwarty metodyki – kwantyfikacja niepewności wejściowych modelu.....	44
3.5.	Element piąty metodyki – walidacji niepewności wejściowych modelu	45
3.6.	Element szósty metodyki – analiza adekwatności oraz ekstrapolacja do obiektu energetyki jądrowej.....	47
4.	Wdrożenie pierwszych elementów metodyki wykonywania obliczeń najlepszego szacowania wraz z oceną niepewności	49
4.1.	Wdrożenie pierwszego elementu metodyki.....	49
4.2.	Wdrożenie drugiego elementu metodyki.....	52
4.3.	Wdrożenie trzeciego elementu metodyki	57
4.4.	Metoda globalnej analizy wrażliwości	59
4.5.	Metody wstecznej kwantyfikacji niepewności	63
4.5.1.	Metoda wstecznej kwantyfikacji niepewności oparta na podejściu Bayesowskim.....	63
4.5.2.	Metoda wstecznej kwantyfikacji niepewności oparta o uczenie maszyn.....	71
5.	Kwantyfikacja niepewności wejściowych na podstawie obliczeń dla instalacji testowych.....	81
5.1.	Obliczenia dla instalacji MARVIKEN	81
5.1.1.	Opis modelowanego zjawiska	81
5.1.2.	Opis instalacji eksperymentalnej Marviken, kwalifikacja modelu (krok 8).....	83
5.1.3.	Wybór parametrów wejściowych (kroki 9 oraz 10)	85

5.1.4.	Kwantyfikacja niepewności parametrów wejściowych (krok 11).....	86
5.1.5.	Sprawdzenie poprawności uzyskanych rozkładów gęstości prawdopodobieństwa (krok 13)	107
5.2.	Obliczenia dla instalacji FLECHT SEASET.....	111
5.2.1.	Opis modelowanych zjawisk.....	112
5.2.2.	Opis instalacji eksperymentalnej FLECHT SEASET, kwalifikacja modelu (krok 8).	113
5.2.3.	Wybór parametrów wejściowych (kroki 9 oraz 10).....	117
5.2.4.	Kwantyfikacja niepewności parametrów wejściowych (krok 11).....	120
5.2.5.	Sprawdzenie poprawności uzyskanych rozkładów gęstości prawdopodobieństwa (krok 13)	143
6.	Walidacja uzyskanych rozkładów gęstości prawdopodobieństwa parametrów – elementy piąty i szóstki metodyki.....	148
6.1.	Obliczenia walidacyjne dla instalacji eksperymentalnej LOFT	148
6.1.1.	Kwalifikacja modelu instalacji LOFT – stan ustalony	151
6.1.2.	Kwalifikacja modelu instalacji LOFT – stan nieustalony, obliczenia referencyjne	153
6.1.3.	Walidacja rozkładów gęstości prawdopodobieństwa parametrów – obliczenia dla instalacji LOFT (element piąty metodyki)	158
6.2.	Obliczenia niepewności dla modelu elektrowni jądrowej typu PWR	162
6.2.1.	Kwalifikacja modelu elektrowni jądrowej – stan ustalony	164
6.2.2.	Kwalifikacja modelu instalacji elektrowni jądrowej – stan nieustalony, obliczenia referencyjne	165
6.2.3.	Obliczenia niepewności stosując uzyskane rozkłady gęstości prawdopodobieństwa parametrów (element szósty metodyki).....	169
7.	Podsumowanie.....	173
7.1.	Ocena metodyki oraz zaproponowanych metod.....	174
7.2.	Wnioski i możliwości kontynuacji badań.....	178
8.	Referencje.....	180
9.	Aneks.....	186
9.1.	Metamodela	186
9.2.	Zestaw testowy – Marviken – wszystkie klasy dla obu zestawów uczących	187
9.3.	Zestaw testowy – Flecht Seaset – wszystkie klasy.....	189

1. Wprowadzenie

1.1. Deterministyczne analizy bezpieczeństwa dla reaktorów jądrowych

Analiza bezpieczeństwa to ocena potencjalnych indywidualnych oraz społecznych zagrożeń związanych z określoną aktywnością, czy jak w przypadku energetyki jądrowej eksploatacją obiektu. W energetyce jądrowej analiza bezpieczeństwa jest częścią składową całościowej oceny bezpieczeństwa obiektu i jest wykonywana w celu weryfikacji czy obiekt jądrowy spełnia wymagania bezpieczeństwa określone przez obowiązujące przepisy i standardy bezpieczeństwa. Zwyczajowo analizy bezpieczeństwa dzieli się na probabilistyczne oraz deterministyczne [1]–[3], które wzajemnie się uzupełniają. Celem probabilistycznych analiz jest w głównej mierze określenie czynników przyczyniających się do powstania zagrożeń radiacyjnych związanych z eksploatacją obiektu jądrowego oraz ocena czy projekt obiektu spełnia probabilistyczne kryteria bezpieczeństwa. Celem deterministycznych analiz bezpieczeństwa jest weryfikacja czy funkcje bezpieczeństwa tj. sterowanie reaktywnością, odprowadzanie ciepła z reaktora oraz osłanianie przed promieniowaniem jonizującym, a także zatrzymywanie substancji promieniotwórczych, ograniczanie i kontrolowanie ich uwolnień do środowiska mogą zostać wypełnione z należytą niezawodnością. Analizy bezpieczeństwa pozwalają również na ocenę czy rozwiązania projektowe obiektu zapewniają właściwą sekwencję poziomów bezpieczeństwa. Dodatkowo weryfikuje się czy obiekt odporny jest na zagrożenia zewnętrzne i wewnętrzne, a struktury, systemy i komponenty oraz działania operatorów pozwalają na utrzymanie uwolnień substancji promieniotwórczych poniżej ustalanych w prawie limitów. Efektem przeprowadzonych analiz powinno być wykazanie z wysokim stopniem pewności, że zagrożenia dla środowiska, pracowników oraz ogółu ludności będą na akceptowalnie niskim poziomie. Osiąga się to poprzez porównanie wyników analiz bezpieczeństwa z kryteriami akceptacji i wymaganiami technicznymi zdefiniowanymi w przepisach prawa [4]–[6] oraz w normach technicznych. Analizy bezpieczeństwa dla danego obiektu rozwijane są wraz z postępem projektu danego obiektu, zaś po jego ukończeniu poddaje się je ciągłemu przeglądowi oraz aktualizacji.

Deterministyczne analizy bezpieczeństwa prowadzi się dla stanów eksploatacyjnych obiektu tj. dla normalnej eksploatacji i przewidywanych zdarzeń eksploatacyjnych oraz dla warunków awaryjnych tj. dla awarii projektowych oraz rozszerzonych warunków projektowych. Analizowane jest działanie rozwiązań projektowych, systemów zabezpieczeń oraz systemów bezpieczeństwa obiektu jądrowego w odpowiedzi na postulowane zdarzenie inicjujące lub kombinację kilku postulowanych zdarzeń inicjujących, które zostały przypisane do każdego ze

stanów obiektu. Analiza warunków normalnej eksploatacji obiektu polega przede wszystkim na weryfikacji czy dawki dla pracowników oraz ogółu ludności otrzymywane w wyniku eksploatacji obiektu są poniżej limitów we wszystkich możliwych spodziewanych trybach pracy obiektu obejmujących przykładowo przeładunek paliwa, stany powyłłączeniowe, uruchamianie reaktora i inne. Przewidywane zdarzenia eksploatacyjne to zdarzenia stanowiące odchylenia od normalnej eksploatacji, które mogą występować podczas pracy reaktora, jednak ich skutki nie powodują zagrożenia dla ludności oraz środowiska. Zdarzenia te powinny być opanowywane poprzez systemy zabezpieczeń reaktora bez konieczności stosowania systemów bezpieczeństwa. Systemy te dedykowane są do opanowywania awarii projektowych. Awarie projektowe to awarie o potencjalnie większych konsekwencjach radiologicznych, w związku z czym powinny mieć niską częstość wystąpienia. W ramach analiz awarii projektowych ocenia się czy obiekt został zaprojektowany tak, żeby skutki radiologiczne awarii projektowych były jak najniższe, reaktor był bezpiecznie wyłączony, a ciepło powyłłączeniowe odbierane. Ostatnią kategorią stanów obiektu są rozszerzone warunki projektowe, w wyniku których może dojść do uwolnień substancji promieniotwórczych do środowiska. Projekt obiektu jądrowego powinien zapewniać ograniczanie skutków takich awarii w celu utrzymania uwolnień substancji promieniotwórczych na najniższym praktycznie możliwym poziomie. Osiąga się to poprzez uwzględnienie w projekcie elektrowni dedykowanych rozwiązań technicznych oraz, w razie konieczności, zastosowanie wdrożonych procedur i planów awaryjnych.

Wykonywanie deterministycznych analiz bezpieczeństwa wspomagane jest poprzez przeprowadzanie obliczeń koniecznych do przewidzenia odpowiedzi obiektu na rozpatrywane zdarzenie inicjujące. Obliczenia te można podzielić na podstawie ich głównego zastosowania na ciepłno-przepływowe, neutronowe, mechaniczne, wytrzymałościowe czy radiologiczne. Wykonuje się je z zastosowaniem specjalistycznych i dedykowanych programów oraz kodów komputerowych. W przypadku obliczeń ciepłno-przepływowych ich wynikami są wartości wielkości fizycznych takich jak temperatura koszulki paliwowej, strumień cieplny, strumień neutronów, ciśnienie i masowe natężenie przepływu w obiegu pierwotnym, koncentracja gazów palnych w obudowie bezpieczeństwa i wiele innych. Parametry te pozwalają na ocenę czy kryteria akceptacji określone dla stanów obiektu zostały spełnione.

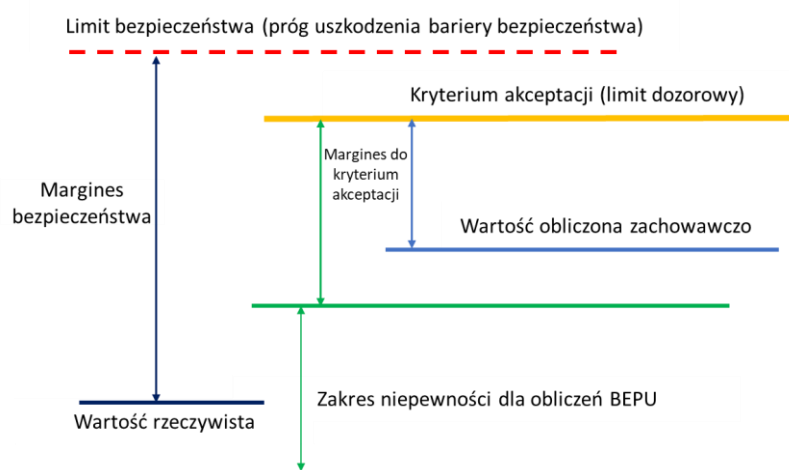
Deterministyczne analizy bezpieczeństwa mogą być wykonywane stosując różne podejścia do założeń dotyczących dostępności systemów obiektu, przewidywanych uszkodzeń elementów wyposażenia, wartości parametrów wejściowych, warunków początkowych i brzegowych. Wyróżnia się dwa główne podejścia: zachowawcze oraz najlepszego szacowania. Historycznie

pierwsze analizy bezpieczeństwa wykonywano stosując przybliżone metody obliczeniowe bądź ograniczone kody obliczeniowe, posiadające wiele zachowawczych założeń modelowych oraz charakteryzujące się małą wydajnością obliczeniową. Stąd przyjęto nazwę tego typu analiz – zachowawcze lub konserwatywne. Podejście zachowawcze cechuje się więc takimi założeniami dotyczącymi parametrów, które powodują niekorzystny z punktu widzenia potencjalnych konsekwencji wpływ tychże parametrów na przebieg scenariusza awaryjnego. Dodatkowo w związku z ograniczonym stanem wiedzy dotyczącej zjawisk zachodzących podczas awarii w obiektach jądrowych stosowano zachowawcze założenia dotyczące warunków początkowych, brzegowych oraz dostępności systemów. W ten sposób uzyskiwano pewność, że nawet dla najbardziej pesymistycznych warunków obiekt będzie zachowywał się bezpiecznie. W praktyce okazało się, że istnieją scenariusze, które mimo zachowawczych założeń, nie prowadzą do najniekorzystniejszych skutków. Obecnie stosuje się bardziej rozwinięte kody obliczeniowe – kody najlepszego szacowania. Jednak przy stosowaniu tego podejścia wiele założeń wciąż pozostaje zachowawczych, tak aby otrzymać wartości parametrów wyjściowych z analizy z wystarczającym marginesem bezpieczeństwa. Według polskich przepisów prawnych podejście zachowawcze powinno być wykorzystywane przy analizie przewidywanych zdarzeń eksploatacyjnych oraz awarii projektowych.

Wraz z rozwojem kodów obliczeniowych oraz wiedzy dotyczącej zjawisk fizycznych zachodzących podczas różnych scenariuszy awaryjnych zaczęto dążyć do przeprowadzania obliczeń oraz analiz bezpieczeństwa opartych o bardziej realistyczne założenia [7], [8]. Rozpoczęto więc opracowywać metody najlepszego szacowania wraz z oceną niepewności (ang. „Best Estimate Plus Uncertainty” – BEPU), które obejmują swoim zakresem przede wszystkim uwzględnienie i szacowanie niepewności, standaryzację opracowywania modelu matematycznego i jego kwalifikację oraz uwzględnienie efektu skali. Metody BEPU pozwalają przede wszystkim na wykorzystanie całej dostępnej obecnie wiedzy bez zbędnego i czasami wprowadzającego w błąd konserwatyzmu, przy zachowaniu zorientowania na weryfikację bezpieczeństwa obiektu. W analizie przyjmuje się bardziej realistyczne założenia dotyczące parametrów wejściowych, warunków początkowych i brzegowych, najczęściej wraz z określeniem ich zakresu niepewności wokół wartości oczekiwanej. W zależności od podejścia do analiz BEPU stosuje się zachowawcze założenia co do dostępności systemów lub realistyczne oparte np. na analizie probabilistycznej bazującej na danych niezawodnościowych. Zastosowanie metod BEPU wymaga zdecydowanie większych zasobów obliczeniowych niż

w przypadku analizy zachowawczej. Metody te wedle polskich wymagań wykorzystuje się w obliczeniach rozszerzonych warunków projektowych.

Wyniki obliczeń wspomagających deterministyczne analizy bezpieczeństwa należy porównać z kryteriami akceptacji określonymi dla danego stanu obiektu. Jedną z najczęściej wykorzystywanych wielkości jest maksymalna temperatura koszulki paliwowej (PCT – ang. „Peak Cladding Temperature”), która pozwala na ocenę czy pierwsza bariera ochronna – koszulka paliwowa pozostała nienaruszona. W przypadku obliczeń zachowawczych wynikiem jest jedna skalarna wartość, którą porównuje się z kryterium. Natomiast dla analiz wykorzystujących metody BEPU jest to wektor wartości, spośród którego można wyróżnić wartości skrajne (minimum i maksimum) stanowiące zakres niepewności oraz wartość średnią stanowiącą wartość najlepszego szacowania. Różnice w określaniu wartości wielkości stanowiących podstawę do oceny kryteriów akceptacji, również w odniesieniu do marginesów bezpieczeństwa przedstawiono na rysunku 1.1.



Rys.1.1. Marginesy bezpieczeństwa oraz porównanie z kryterium akceptacji w analizach zachowawczych oraz BEPU [9]

Analizy wykonywane stosując podejście zachowawcze skutkują uzyskaniem wyników odbiegających od wartości rzeczywistej, ale wciąż poniżej wartości wyznaczonej przez kryterium akceptacji. Wynika to wprost z założeń i celu wykonywania takich analiz, w konsekwencji przyjmuje się, że margines do kryterium akceptacji może być niewielki. Należy jednak zwrócić uwagę, że kryterium akceptacji zakładane jest z pewnym zapasem w odniesieniu do limitów bezpieczeństwa, które stanowią wartości progowe. Analizy BEPU wykorzystując bardziej realistyczne założenia powinny pozwalać na uzyskiwanie rezultatów obejmujących wartość rzeczywistą, przy jednoczesnym zapewnieniu większego marginesu bezpieczeństwa niż obliczenia zachowawcze. Na rysunku 1.1 jest to schematycznie

zaprezentowane w taki sposób, że wartość rzeczywista mieści się w zakresie niepewności uzyskanym dla obliczeń BEPU. Margines do kryterium akceptacji wyznaczany od górnej granicy zakresu niepewności jest, również większy niż w obliczeniach zachowawczych. Różnice te ukazują rozwój analiz bezpieczeństwa i ich orientację na uwzględnianie bardziej realistycznych założeń i uzyskiwanie bardziej wiarygodnych marginesów bezpieczeństwa.

Obecnie stosowane są dwie główne metody najlepszego szacowania wraz z oceną niepewności: propagacji niepewności wejściowych oraz propagacji niepewności wyjściowych. W środowisku badawczym analityków systemów reaktorów jądrowych opracowywane są obecnie nowe metody wykorzystujące najlepsze cechy obu głównych metod (propagacji niepewności wejściowych oraz wyjściowych) w połączeniu z analizami wrażliwości. Wprowadzają one dodatkowe kroki w metodykach BEPU, poprzez dodatkową wsteczną kwantyfikację niepewności (ang. „Inverse Uncertainty Quantification” – IUQ).

1.2. Zagadnienie niepewności w deterministycznych analizach bezpieczeństwa

Niepewność w naukach fizycznych oznacza brak częściowej bądź całkowitej wiedzy na temat zjawiska czy informacji. Niepewności mogą być dwojakiego rodzaju: statystycznego (po angielsku określane jako „aleatoric uncertainties”) oraz epistemicznego, czyli systematycznego. Niepewności statystycznych nie da się zredukować do zera. Są to niepewności wprost związane ze stochastycznością procesów. Niepewności epistemiczne można próbować redukować poprzez zwiększanie wiedzy na temat rozkładu wartości źródła niepewności epistemicznej (np. wartości wielkości fizycznej czy parametru modelu). W modelowaniu matematycznym większość niepewności wynikających z niedoskonałości modelu w odwzorowaniu rzeczywistości jest epistemiczna, co oznacza, że istnieje pole do ich redukcji. Nie ma jednak możliwości ich całkowitej redukcji, ponieważ zawsze pozostaje niezerowa wariancja badanej wartości [10].

Główne źródła niepewności w analizach bezpieczeństwa stanowią [11]–[13]:

- Niepewności modeli i kodów obliczeniowych – kody obliczeniowe wykorzystywane w obliczeniach dla analiz bezpieczeństwa oparte są o modele fizyczne które już w punkcie wyjścia są niedoskonałym odwzorowaniem rzeczywistości. Należy uwzględnić niedoskonałości stosowanych równań bilansu, wynikające między innymi z konieczności uśrednienia w obliczeniach wielu parametrów fizycznych w taki sposób by opisywały stan komórek obliczeniowych, braku modeli przepływu dwufazowego w pełni odwzorowujących rzeczywistość czy niedokładności równań konstytutywnych,

często uzyskiwanych empirycznie. Równania domykające (jak na przykład modele wymiany ciepła na ściankach) są również często stosowane poza zbadanym zakresem ich poprawności, co wprowadza dodatkowe źródło niepewności. Często stosowane są również dodatkowe specjalne modele (uzyskane empirycznie) aby odwzorować zjawiska zachodzące w czasie awarii w elektrowniach jądrowych (przykładowo dla zjawiska przepływu przeciw-prądowego). Ze względu na stosowanie stabelaryzowanych bądź obliczonych z wykorzystaniem wielomianów (zawierających empirycznie wyznaczone stałe) wartości takich jak ciepło właściwe, przewodność cieplna czy gęstość dla danego materiału to właściwości materiałowe stosowane w kodach obliczeniowych stanowią kolejne źródło niepewności. Ostatnim istotnym czynnikiem jest stosowanie przybliżonych metod numerycznych do rozwiązywania cząstkowych równań różniczkowych, co również wprowadza niemierzalne niepewności. W niektórych opracowaniach wydziela się kolejną grupę niepewności związanych z efektem skalowania, jednak jest to de facto podgrupa niepewności modeli i kodów obliczeniowych. Wiele równań domykających oraz specjalnych modeli zostało opracowanych na podstawie rezultatów przeprowadzonych eksperymentów, które z zasady nie mogły być przeprowadzone w rzeczywistych warunkach eksploatacyjnych obiektów jądrowych. Instalacje eksperymentalne budowane były więc w pewnej skali i mogły odwzorowywać jedynie wybrane zjawiska lub wiele zjawisk zachodzących podczas określonego scenariusza awarii. Nie można jednak zakładać ze stuprocentową pewnością, że odpowiedź systemów obiektu energetyki jądrowej będzie identyczna jak instalacji testowej w mniejszej skali.

- Niepewności nodalizacji – mimo istnienia instrukcji dla użytkowników kodów nie istnieją jedne, klarowne wytyczne opracowywania tzw. nodalizacji (czyli schematu reprezentacji obiektu w modelu, poprzez wyznaczanie liczby oraz geometrii i połączeń komórek – „nodów” każdego komponentu składającego się na modelowany system). Wielu użytkowników może więc opracować wiele różnych nodalizacji, mimo posiadania tych samych danych wejściowych. Dodatkowo ze względu na specyfikę kodów obliczeniowych nawet małe różnice w nodalizacji mogą wprowadzać znaczące różnice w wartościach wyjściowych.
- Efekt użytkownika – powiązany jest z poprzednim punktem, ponieważ wszelkie decyzje dotyczące nodalizacji, wykorzystania opcji dodatkowych modeli dostępnych w kodzie, wyboru wersji kodu, interpretacji danych czy kwalifikacji wyników stanu ustalonego

pozostają w gestii użytkownika kodu obliczeniowego. Dwie grupy analityków stosując ten sam kod obliczeniowy mogą, więc otrzymać rozbieżne rezultaty.

- Niepewności związane z warunkami w obiekcie – obejmuje to zarówno warunki początkowe i brzegowe jak i założenia dotyczące charakterystyk pomp, oporów hydraulicznych czy współczynników tarcia.

W celu ograniczania pierwszego ze źródeł niepewności stosuje się metody weryfikacji oraz walidacji kodów obliczeniowych. Weryfikacja kodu przeprowadzana jest przez jego dostawcę. Weryfikacja polega na sprawdzeniu czy stosowane w kodzie obliczeniowym modele matematyczne są poprawnie w nim zaimplementowane i rozwiązywane. Uwzględnia to zarówno równania jak i metody numeryczne. Walidacja może być przeprowadzona przez dostawcę kodu jak również przez jego użytkowników. Ma ona na celu określenie czy modele matematyczne stosowane w kodzie w wystarczający sposób odzwierciedlają rzeczywistość. Jest to więc sprawdzenie czy rozwiązywane są poprawne równania w odniesieniu do systemu bądź zjawiska, które jest modelowane. Najlepszą metodą przeprowadzenia walidacji kodu oraz przy okazji kwalifikacji nodalizacji jest porównanie obliczeń uzyskanych z wykorzystaniem kodu obliczeniowego z danymi eksperymentalnymi. Dla określonych scenariuszy awarii oraz zjawisk fizycznych zachodzących podczas takich awarii to porównanie jest wykonywane na podstawie danych z instalacji eksperymentalnych. Instalacje te dzieli się na dwie główne kategorie:

- Integral Test Facilities (ITF) w których badane jest wiele zjawisk zachodzących podczas jednego scenariusza awaryjnego [14],
- Separate Effect Test Facilities (SETF) w których badane są pojedyncze zjawiska w obiektach o mniejszej skali i prostszej konfiguracji eksperymentalnej [15].

Eksperymenty wykonywane w takich instalacjach prowadzone są w kontrolowanym środowisku, z odpowiednią aparaturą pomiarową rejestrującą zmiany parametrów wyjściowych istotnych z punktu widzenia analizy bezpieczeństwa. Posiadając więc dane eksperymentalne można określić dokładność ich odwzorowania przez kody obliczeniowe oraz dokonać korekt nodalizacji i założeń dotyczących wartości wewnętrznych parametrów modeli. Należy mieć jednak wciąż na względzie kolejne trzy źródła niepewności opisane powyżej, szczególnie efekt użytkownika, jak również błąd pomiarowy. W celu eliminacji jak największej liczby źródeł niepewności przeprowadza się proces kwalifikacji obliczeń wykonywanych dla

instalacji doświadczalnej. Proces kwalifikacji opisany jest w [rozdziale 5](#), wraz z demonstracją jego praktycznego zastosowania.

1.3. Motywacja rozprawy

Rozprawa dotyczy zagadnień związanych z energetyką jądrową, przede wszystkim z bezpieczeństwem reaktorów jądrowych. Podstawową motywacją jest opracowanie nowej metodyki wykonywania i weryfikacji deterministycznych analiz bezpieczeństwa z wykorzystaniem podejścia najlepszego szacowania wraz z oceną niepewności. Rozprawa powstała w ramach programu „doktorat wdrożeniowy”, stąd metodyka będzie wdrożona w pracach Państwowej Agencji Atomistyki (PAA), szczególnie w perspektywie planowanego Programu polskiej energetyki jądrowej, gdzie jednym z głównych zadań PAA będzie dokonanie oceny Raportu Bezpieczeństwa obiektu jądrowego wraz z towarzyszącą dokumentacją techniczną. Niezależna weryfikacja przedstawionych wyników obliczeń oraz analiz na podstawie zaproponowanych przez inwestora metodologii jest jednym z kluczowych zadań PAA w procesie wydawania zezwolenia na budowę, rozruch czy eksploatację obiektu jądrowego. Państwowa Agencja Atomistyki powinna więc posiadać zarówno wiedzę i doświadczenie jak i swoje niezależne metody pozwalające na wykonanie obliczeń sprawdzających.

Obecnie stosowane metody wykonywania obliczeń BEPU posiadają niedoskonałości, które środowiska akademickie oraz dozоровe próbują wyeliminować poprzez integrację najlepszych cech tych metod oraz wsparcie się metodami analizy Bayesowskiej i globalnej analizy wrażliwości. W metodach opartych na propagacji niepewności wejściowych szczególnym mankamentem jest duży wpływ użytkownika oraz brak powtarzalności uzyskiwanych wyników. W przypadku drugiej metody opartej o ekstrapolację niepewności największą przeszkodą jest brak uniwersalności tej metody (gdy brak jest danych doświadczalnych dotyczących danego zjawiska metoda nie może zostać wykorzystana) oraz konieczność posiadania bardzo rozbudowanej bazy danych instalacji eksperymentalnych. Obszerniejszy krytyczny opis właściwości obecnych metod znajduje się w [rozdziale 2](#).

Wraz z rozwojem nowej metodyki konieczne jest zaimplementowanie nowych metod obliczeniowych wstecznej kwantyfikacji niepewności parametrów wejściowych. Metody te zostały zaproponowane stosując dwa różne podejścia tak aby porównać ich skuteczność oraz intuicyjność implementacji. Metodyka wraz z metodami obliczeniowymi powinna być uniwersalna, pozwalając na zarówno przeprowadzenie obliczeń sprawdzających dla wybranych scenariuszy awaryjnych jak również dla weryfikacji przyjmowanych wartości parametrów do

analizy wraz z ich rozkładami gęstości prawdopodobieństwa. Metodyka powstała na bazie raportu z projektu SAPIUM [16], którego wyniki wciąż nie są udostępnione publicznie, a sama publikacja ma formę wytycznych oraz przedstawienia dobrych praktyk. Metodyka tego typu nie została więc jeszcze w całości wdrożona w praktyce, a na wybranych jej krokach użytkownik zobowiązany jest sam opracować metody obliczeniowe. W ramach rozprawy wdrożono metodykę z pewnymi modyfikacjami w porównaniu z raportem z projektu SAPIUM i wraz z propozycją własnych metod obliczeniowych dokonano jej pierwszej walidacji w ograniczonym do dwóch zjawisk fizycznych zakresie.

1.4. Cel i zakres rozprawy

Głównym celem rozprawy jest usprawnienie obecnie istniejących metod wykonywania analiz najlepszego szacowania wraz z oceną niepewności oraz wdrożenie usprawnionej metodyki. Usprawniona metodyka powinna korzystać z najlepszych światowych praktyk oraz doświadczeń w stosowaniu obecnych metod, przy jednoczesnym dążeniu do minimalizacji ich wad. Nowa metodyka powinna w pierwszej kolejności pozwolić na minimalizację wpływu użytkownika na określanie wartości i rozkładów gęstości prawdopodobieństwa parametrów opisujących modele matematyczne stosowane w kodach obliczeniowych. Parametry te często określane mianem parametrów kalibracyjnych, mogą w znaczący sposób wpływać na uzyskiwane rezultaty, a tym samym na spełnienie dozorowych kryteriów akceptacji. Dodatkowo istotne jest, aby metodyka była uniwersalna a wyniki obliczeń najlepszego szacowania wraz z oceną niepewności uzyskiwane w wyniku jej stosowania były powtarzalne. Kolejnym celem powinno być sprawdzenie metodyki oraz przeprowadzenie jej weryfikacji na konkretnym przypadku scenariusza awaryjnego w elektrowni jądrowej.

Proponowana teza rozprawy przyjmuje następujące brzmienie:

Istnieją usprawnienia do obecnie stosowanej metody propagacji niepewności wejściowych, które pozwolą na ograniczenie wpływu użytkownika oraz zapewnienie uniwersalności oraz powtarzalności metody.

Bezpośrednią metodą udowodnienia tezy będzie opracowanie nowej metodyki oraz wdrożenie jej z zastosowaniem nowych metod obliczeniowych. Metoda powinna bazować na obecnych doświadczeniach, danych eksperymentalnych oraz być zweryfikowana w odniesieniu zarówno do danych eksperymentalnych jak i przewidywanych scenariuszy awaryjnych w elektrowniach jądrowych.

Zasadniczo opracowana metodyka powinna być uniwersalna i możliwa do wykorzystania w analizie wszelkich scenariuszy awaryjnych. Wdrożenie i weryfikacja metodyki dla wielu awarii jest jednak zadaniem, które wymaga zazwyczaj pracy całego zespołu analitycznego bądź wręcz przeprowadzenia międzynarodowego benchmarku. Zakres opracowania, wdrożenia i weryfikacji metodyki został więc ograniczony do trzech podstawowych zjawisk badanych podczas awarii dużego rozerwania rurociągu obiegu chłodzenia (LBLOCA – ang „Large Break Loss of Coolant Accident”). Weryfikację przedstawiono dla tejże awarii w reaktorze typu PWR, jednak sama metodyka powinna być niezależna od technologii analizowanego reaktora.

2. Metody obliczeń najlepszego szacowania wraz z oceną niepewności - stan wiedzy

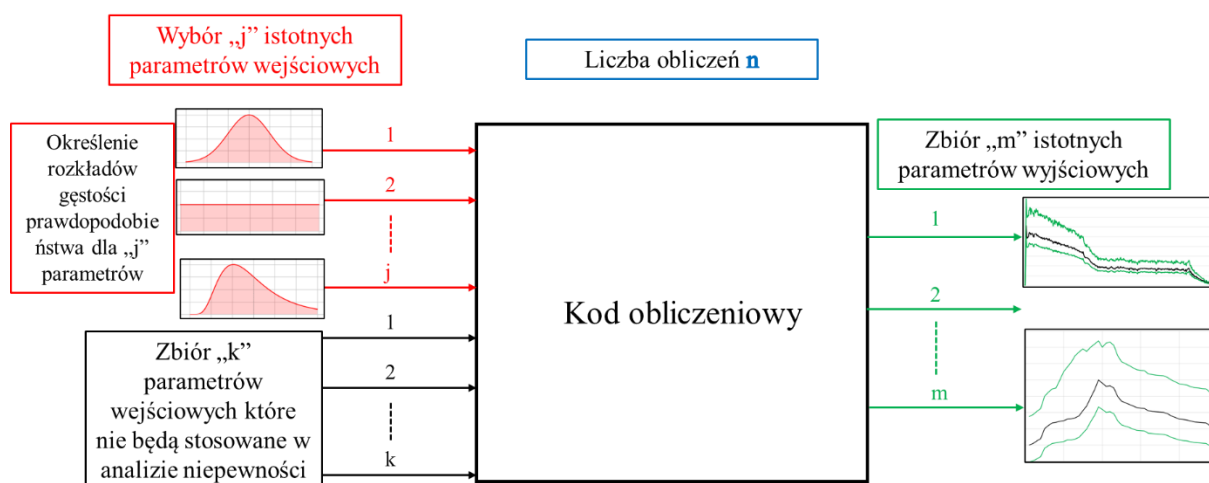
Pierwsze podejścia do stworzenia metod wykonywania obliczeń najlepszego szacowania wraz z oceną niepewności pojawiały się w Stanach Zjednoczonych w latach osiemdziesiątych dwudziestego wieku. Było to wynikiem znacznego zwiększenia nakładów finansowych na badania dotyczące zjawisk ciepłno-przepływowych zachodzących w trakcie awarii z rozerwaniem rurociągu obiegu pierwotnego po awarii w elektrowni jądrowej Three Mile Island - 2. Rozwój wiedzy na temat zjawisk fizycznych oraz przebiegu awarii prowadził do wykonywania analiz i obliczeń w których można było unikać zachowawczych założeń. Łącząc to z obserwowalnym w latach osiemdziesiątych rozwojem kodów obliczeniowych, powstawała konieczność szacowania niepewności związanych z uzyskiwanymi rezultatami. Doprowadziło to do opracowania metodyki CSAU [17] (ang. „Code Scaling, Applicability and Uncertainty”), która stanowiła zbiór wytycznych i kładła duży nacisk na osąd ekspercki w ramach procesu identyfikacji kluczowych dla danej awarii zjawisk fizycznych. W ramach tej metodyki nie proponowano jednak żadnych metod obliczeniowych, wskazywano raczej kolejne kroki prowadzące do wykonania analizy niepewności. Następnie w niektórych europejskich ośrodkach badawczych oraz jednostkach pracujących dla operatorów elektrowni jądrowych posiadających kompetencje w analizach deterministycznych dla potrzeb energetyki jądrowej zaczęto opracowywać swoje własne metody. Prace prowadzono w Niemczech w Instytucie Gesellschaft für Anlagen- und Reaktorsicherheit [18] (GRS), w Wielkiej Brytanii w jednostce pracującej dla operatora elektrowni jądrowej w Winfrith (Atomic Energy Authority Winfrith – AEA), w Hiszpanii w Empresa Nacional del Uranio (ENUSA), we Francji w Institut de radioprotection et de sûreté nucléaire (IRSN) czy we Włoszech na Uniwersytecie w Pizie. Dokładny opis każdej z tych metod można odnaleźć w literaturze [9], [19], jednak metody te można podzielić na dwie główne grupy: propagacji niepewności wejściowych oraz propagacji niepewności wyjściowych. Metody BEPU pozwalają na zwiększenie marginesów bezpieczeństwa, stąd operatorzy elektrowni jądrowych mogli udowadniać bezpieczeństwo swoich bloków przy pracy ze zwiększoną mocą. Ze względu na idące za tym korzyści ekonomiczne jakie operatorzy elektrowni jądrowych mogli osiągnąć dzięki stosowaniu metod BEPU, od lat 90 trwały prace nad pierwszymi zastosowaniami metod BEPU (Best Estimate Plus Uncertainty) w analizach wykorzystywanych w licencjonowaniu obiektów jądrowych [20], [21]. Pierwszym przypadkiem było zastosowanie tej metody dla analiz elektrowni jądrowej Angra-2 w Brazylii [22]. Kolejnymi krajami które licencjonowały nowe elektrownie

bądź zezwalały na przedłużenie eksploatacji reaktorów jądrowych wykorzystując metody BEPU były USA, Niderlandy, Niemcy, Litwa czy Korea Południowa.

2.1. Obecnie stosowane metody najlepszego szacowania wraz z oceną niepewności w deterministycznych analizach bezpieczeństwa

2.1.1. Metoda propagacji niepewności wejściowych

Metoda propagacji niepewności wejściowych jest najpowszechniejszą obecnie metodą BEPU, rozwijaną przez wiele ośrodków badawczych, wykorzystywaną w praktyce przy licencjonowaniu reaktorów oraz zaakceptowaną przez wybrane zagraniczne dozory jądrowe. Schemat poglądowy metody przedstawiono na rysunku 2.1. Wartości wybranych parametrów losowane są z określonych rozkładów gęstości prawdopodobieństwa i wprowadzane do pliku wsadowego dla wybranego kodu obliczeniowego. Następnie wykonuje się „n” obliczeń z zastosowaniem tego kodu. Wynikiem są zakresy niepewności określonych wielkości wyjściowych.



Rys. 2.1. Schemat metody propagacji niepewności wejściowych

Metoda ta pozwala na uwzględnienie niepewności wielu parametrów wejściowych, ponieważ liczba koniecznych do wykonania obliczeń jest niezależna od liczby rozpatrywanych parametrów wejściowych. W praktyce, mimo że liczba parametrów które można by uwzględnić w analizie liczona jest w setkach, ogranicza się wybór parametrów do kilkudziesięciu (na Rys.2.1. są to parametry od 1 do „j” zaznaczone na czerwono). Wybór parametrów powinien zostać poprzedzony analizą jakie zjawiska wystąpią podczas scenariusza, dla którego będą wykonywane obliczenia, krok ten jest już więc zależny od eksperckiej wiedzy analityków. Obecnie można korzystać z matryc identyfikujących takie zjawiska dla danych scenariuszy awarii, które zespoły ekspertów międzynarodowych opracowywały pod egidą OECD[23]. Po określeniu listy parametrów dokonywane jest określenie ich zakresów oraz rozkładów gęstości

prawdopodobieństwa. Ten krok charakteryzowany jest silnym wpływem użytkownika, ponieważ brak jest jednolitych wytycznych dotyczących określania rozkładów gęstości prawdopodobieństwa, szczególnie dla parametrów modelowych. Kolejnym krokiem jest określenie liczby obliczeń (iteracji) koniecznych do osiągnięcia zakładanego poziomu tolerancji. Metoda opracowana przez GRS rozpropagowała stosowanie w tym celu nierówności (formuły) Wilksa [24]. Nierówności te pozwalają na wyznaczenie minimalnej liczby obliczeń dla oczekiwanego zakresu tolerancji. Metody propagacji niepewności wejściowych są analizami statystycznymi, stąd konieczne jest ustalenie oczekiwanych przedziałów tolerancji dla rezultatów. Przedział tolerancji jest to region statystyczny wyznaczony przez oczekiwany poziom prawdopodobieństwa uzyskania wyniku spełniającego kryteria akceptacji przy jednoczesnym spełnieniu określonego poziomu pewności rezultatu statystycznego, zgodnie z nierównością:

$$P_S \{ P_Y \{ Y < L \} \geq \alpha \} \geq \beta \quad (2.1)$$

gdzie:

P_S - to prawdopodobieństwo statystyczne, wynikające z nieuniknionej niepewności związanej z wykonywaniem skończonej liczby obliczeń,

P_Y - to prawdopodobieństwo uzyskania wielkości wyjściowej Y , związane z niepewnościami propagowanymi przez kod obliczeniowy

Y - to zmienna losowa analizowanej wielkości wyjściowej

L - jest to dozоровe kryterium akceptacji dla danej wielkości wyjściowej (np. 1200 °C dla temperatury koszulki paliwowej)

α - to ustalony poziom pokrycia (prawdopodobieństwa) – zazwyczaj wynosi 95%

β - to ustalony poziom ufności – zazwyczaj wynosi 95%

Poziom pokrycia można zweryfikować poprzez odnalezienie dziewięćdziesiąt piątego percentyla rozkładu prawdopodobieństwa wielkości wyjściowej. Najczęściej stosowany przedział tolerancji 95/95 oznacza więc, że z dziewięćdziesięciopięcioprocentową pewnością uzyskuje się wartość dziewięćdziesiąt-piętego percentyla mniejszą od wartości kryterium akceptacji. W tabeli 2.1 przedstawiono wymagane liczby obliczeń dla jednostronnego przedziału tolerancji 95/95 w zależności od stopnia formuły Wilksa. Stopień formuły wyznacza która z uporządkowanych rosnąco wartości parametru wyjściowego stanowi percentyl wymaganego przez dozory jądrowe rzędu. W praktyce oznacza to, że w przypadku np. trzeciego

stopnia formuły można wykonać 124 obliczenia i do analizy porównawczej z kryterium akceptacji wybrać sto dwudziestą drugą najwyższą wartość. Wyższe stopnie formuły pozwalają użytkownikowi odrzucać wybraną liczbę nieudanych obliczeń wykonanych z zastosowaniem kodu komputerowego.

Tabela 2.1. Minimalna liczba obliczeń dla przedziału tolerancji 95/95 w zależności od stopnia formuły Wilksa

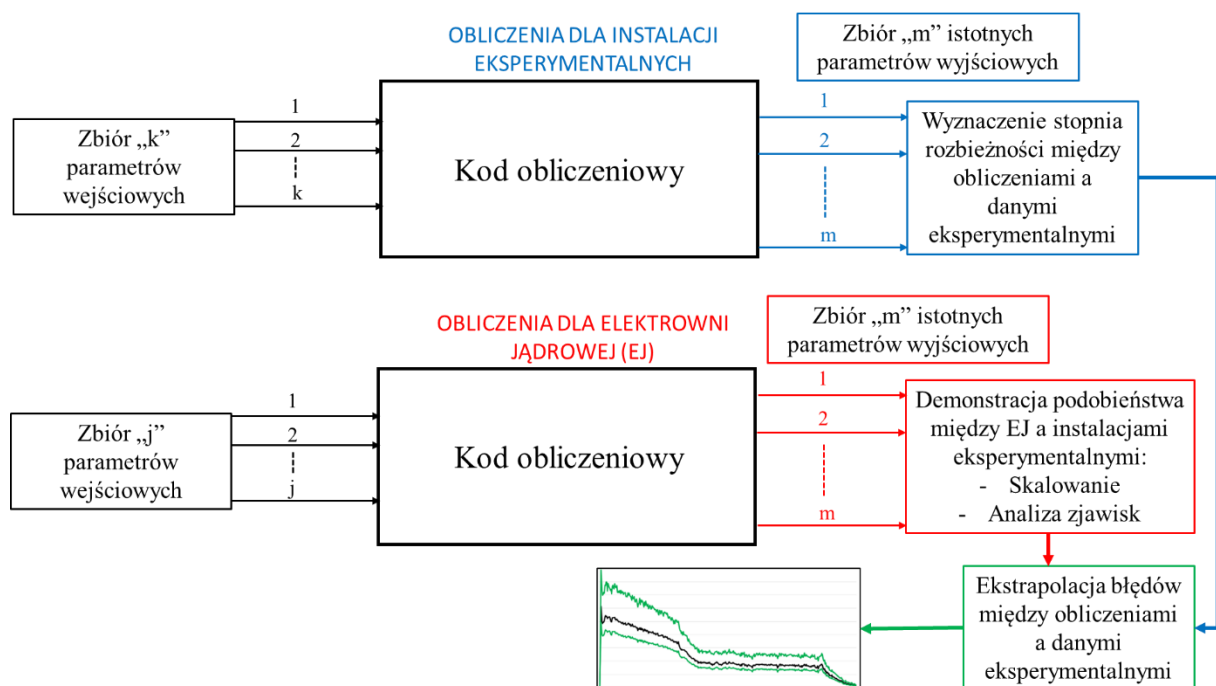
Liczba obliczeń – „n”	Stopień formuły	Wartość stanowiąca 95 percentyl w jednostronnym przedziale tolerancji
59	1	Y (n)
93	2	Y (n-1)
124	3	Y (n-2)
153	4	Y (n-3)
181	5	Y (n-4)

Po wyznaczeniu liczby obliczeń należy wykonać je z zastosowaniem kodu komputerowego. Przygotowując pliki wsadowe dla kodu obliczeniowego dla każdego z wyselekcjonowanych istotnych parametrów wejściowych losuje się wartość z rozkładu gęstości prawdopodobieństwa określonego dla tego parametru. Wartości pozostałych parametrów (zbiór „k” na rysunku 2.1) pozostają takie same we wszystkich obliczeniach. Rezultatem każdego z obliczeń jest zestaw wartości skalarnych bądź przebiegów parametrów istotnych z punktu widzenia kryteriów akceptacji. Dla „n” wykonanych obliczeń tworzą się więc zbiory „n” przebiegów czy „n” wartości skalarnych. Spośród tych zbiorów można wyznaczyć wartości minimalne, maksymalne oraz średnie jak schematycznie zaprezentowano na rysunku 2.1. Zgodnie z opisem dotyczącym zastosowania formuły Wilksa porównuje się następnie wartość stanowiącą np. 95 percentyl z wartością kryterium akceptacji. W następnych krokach analizy można dokonać analizy istotności wpływu permutowanych parametrów wejściowych na wyniki obliczeń.

2.1.2. Metoda propagacji niepewności wyjściowych

Istnieje obecnie jedynie jedna opisana i sprawdzona w praktyce metoda propagacji niepewności wyjściowych – Metoda UMAE (ang. „Uncertainty methodology based on accuracy extrapolation”) [25]. Schemat tej metody przedstawiono na rysunku 2.2. Opiera się ona na ekstrapolacji miary różnicy między danymi eksperymentalnymi oraz wynikami obliczeń dla instalacji eksperymentalnych. Dane z instalacji eksperymentalnych muszą być dopasowane stricte pod analizowany scenariusz, tak aby odwzorować wszystkie zjawiska fizyczne które są oczekiwane do wystąpienia w trakcie danego scenariusza. Konieczne jest, więc posiadanie obszernej bazy eksperymentów w większej skali (ITF), do tego zróżnicowanych pod względem analizowanych scenariuszy awaryjnych. Dla każdej instalacji eksperymentalnej należy więc

wykonać obliczenia z realistycznymi założeniami dotyczącymi warunków początkowych i brzegowych, przeprowadzić proces kwalifikacji nodalizacji, wyników stanu ustalonego oraz nieustalonego, tak aby uzyskać pewność, że stworzony model odpowiednio odwzorowuje eksperyment. Dopiero gdy dowiedzie się poprawności opracowanej nodalizacji można przeprowadzić analizę rozbieżności między danymi eksperymentalnymi oraz obliczeniami wykonanymi z zastosowaniem kodu obliczeniowego. Dla wszystkich eksperymentów wykorzystywanych w ekstrapolacji niepewności oraz modelu elektrowni jądrowych powinna zostać zastosowana nodalizacja opracowana według tej samej metody bądź procedury. Obliczenia dla elektrowni jądrowej wykonuje się jednokrotnie, po czym należy przeprowadzić analizę podobieństwa oraz skalowania, czyli weryfikacji czy wszelkie zjawiska fizyczne obserwowane w mniejszej skali (eksperymentalnej) będą miały taki sam przebieg w skali większego obiektu energetyki jądrowej (jak elektrownia jądrowa). Końcowe zakresy niepewności dla obliczeń dla elektrowni jądrowej wynikają, więc z ekstrapolacji różnic między danymi eksperymentalnymi i obliczeniami dla tych instalacji eksperymentalnych z których pobierano dane eksperymentalne i w których badano przebieg analizowanej awarii.



Rys. 2.2. Schemat metody propagacji niepewności wyjściowych

2.1.3. Metody wykorzystujące analizę Bayesowską oraz globalne analizy wrażliwości

Obecnie podejmowane są próby połączenia najlepszych cech obu opisanych metod, tak aby jednocześnie minimalizować ich wady. W metodach łączących oba podejścia niepewności parametrów wejściowych są wciąż propagowane przez kod obliczeniowy, jednak wykorzystuje

się obliczenia dla instalacji eksperymentalnych w celu określenia rozkładów gęstości prawdopodobieństwa dla tych parametrów. Proces ten jest nazywany wsteczną kwantyfikacją niepewności (ang. „Inverse Uncertainty Quantification” – IUQ) i ma podstawy w analizie Bayesowskiej [26]. Dużą popularność zyskują metody Bayesowskie wspomagane poprzez wykonywanie obliczeń Monte Carlo [27]–[31]. Stosując podejście Bayesowskie zakłada się rozkłady gęstości prawdopodobieństwa a’priori, a następnie poprzez obliczenia funkcji wiarygodności dla obliczeń ze zmieniającymi wartościami badanych parametrów dokonuje się aktualizacji rozkładów. Rozwiązaniem problemu wstecznej kwantyfikacji w podejściu Bayesowskim są więc rozkłady gęstości prawdopodobieństwa uzyskane a’posteriori. Szczegółowe omówienie zastosowania podejścia Bayesowskiego do wstecznej kwantyfikacji niepewności znajduje się w [rozdziale 4.5.1](#). Metody Monte Carlo są bardzo wymagające obliczeniowo stąd często wprowadzane są metamodele [32], czyli funkcje stosowane do aproksymacji relacji między danymi wejściowymi a wyjściowymi modelu matematycznego który zastępują. W przypadku analiz deterministycznych metamodele zastępują więc kod obliczeniowy. Główny koszt obliczeniowy przenoszony jest na wytrenowanie metamodelu. Otwarta pozostaje jednak kwestia dokładności metamodeli w porównaniu do zastosowania kodu obliczeniowego. Przykłady obecnie stosowanych metod wykorzystujących podejście Bayesowskie pokrótce przedstawiono poniżej.

Metoda CIRCE [33] została opracowana w efekcie współpracy specjalistów z EDF (Électricité de France), Arevy, IRSN oraz CEA (Commissariat à l’énergie atomique). Metoda ta początkowo dostosowana była jedynie pod kod ciepłno-przepływowy CATHARE, jednak obecnie można ją stosować również do innych kodów obliczeniowych. CIRCE pozwala na estymację wartości średniej, odchylenia standardowego oraz rozkładu gęstości prawdopodobieństwa, spośród dwóch propozycji – rozkładu normalnego bądź logarytmiczno-normalnego dla badanych parametrów. Metoda wykorzystuje podejście Bayesowskie oraz estymację metodą największej wiarygodności. Jednym z założeń metody jest liniowość zależności między badanym parametrem wejściowym, a wielkościami reprezentującymi wyniki obliczeń kodu. Kolejnym założeniem jest brak zależności badanych parametrów wejściowych od siebie, co pozwala na uproszczenie analizy poprzez eliminację członów reprezentujących kowariancję między parametrami.

Podobne założenia są wykorzystywane w metodzie MCDA opracowanej przez KAERI (Korea Atomic Energy Research Institute) [34]. W metodzie tej wykorzystywane są metody deterministyczne oparte o obliczenia szeregów Taylora dla wielkości reprezentujących wyniki

obliczeń kodu. W przypadku braku liniowości między parametrami wejściowymi i wyjściowymi, w metodzie wykorzystywane są metody Markov Chain Monte Carlo.

Poza metodami opartymi na obliczeniach Monte Carlo, opracowywano również inne metody między innymi metodę FFTBM, która była wykorzystywana przez niektórych uczestników Benchmarku PREMIUM opisanego w kolejnym podrozdziale. Na tej podstawie rozwinięto metodę IPREM [35] opierającą się o wykonywanie lokalnych obliczeń wrażliwości (a więc dla każdego z parametrów wejściowych oddzielnie) i zastosowanie szybkiej transformacji Fouriera do porównania wyników obliczeń i danych eksperymentalnych. Rozkłady proponowane przez tę metodę przyjmują formę rozkładów jednostajnych, metoda pozwala więc przede wszystkim na określenie zakresów zmienności parametrów.

Pierwsze próby usystematyzowania wiedzy oraz nowych metod i podejść rozwijanych w różnych ośrodkach badawczych podjęto w ramach, opisanego w [rozdziale 2.2.2](#), projektu PREMIUM pod przewodnictwem NEA (Nuclear Energy Agency) OECD [36], [37].

2.2. Ograniczenia obecnie stosowanych metod

Obie opisane powyżej najbardziej rozpowszechnione metody najlepszego szacowania wraz z oceną niepewności posiadają wady, które środowisko analityków bezpieczeństwa obiektów jądrowych stara się minimalizować. W przypadku metody propagacji niepewności wejściowych, do głównych jej wad można zaliczyć znaczne uzależnienie od wpływu użytkownika, jego decyzji podejmowanych na etapie opracowywania modelu oraz brak powtarzalności. Efekt użytkownika jest obserwowalny w przeprowadzeniu procesu minimalizacji liczby parametrów wejściowych i określaniu zakresu zmienności oraz rozkładu gęstości prawdopodobieństwa dla wybranych parametrów. Dodatkowo efekt użytkownika obserwowany jest również każdorazowo przy opracowywaniu nodalizacji, a więc co za tym idzie przy wykonywaniu obliczeń dla opracowanej nodalizacji. Brak powtarzalności wynika stricte z ograniczonej liczby wykonywanych obliczeń, czyli ograniczonej próby statystycznej w odniesieniu do permutowanych parametrów wejściowych. Liczba wykonywanych obliczeń wynosi zazwyczaj około 100. Wykonując kilkakrotnie zestaw stu obliczeń, za każdym razem losowo próbując wartości z rozkładów niepewności parametrów wejściowych, ze względu na małą próbę uzyska się inne zestawy wartości tych parametrów, a tym samym rozbieżne będą wartości badanych wielkości wyjściowych. Problem ten jeszcze szczególnie istotny dla parametrów, dla których użytkownik posiada niepełną wiedzę na temat ich wartości oczekiwanych i rozkładów, w efekcie użytkownik powinien założyć szeroki zakres

niepewności tych parametrów, co przy małej próbie doprowadzi do małej powtarzalności w losowaniu wartości.

Metoda propagacji niepewności wyjściowych wymaga posiadania rozległej bazy danych doświadczalnych z instalacji eksperymentalnych dla których następnie należy opracować nodalizacje, modele matematyczne z wykorzystaniem kodu obliczeniowego, dokonać kwalifikacji nodalizacji oraz wykonać obliczenia, które porówna się z danymi eksperymentalnymi. Proces ten jest bardzo czasochłonny, stąd opracowanie odpowiedniej bazy eksperymentów do stosowania metody jest długotrwały. Dodatkowo, ze względu na fakt, że dokonuje się ekstrapolacji niepewności wyjściowych, a więc nie ma bezpośredniej relacji między niepewnościami wejściowymi a wyjściowymi to w metodzie tej nie można przeprowadzić analizy wpływu parametrów wejściowych na badane parametry wyjściowe. Znaczącym ograniczeniem jest również fakt, że jeśli dla wybranego scenariusza awaryjnego nie istnieje wystarczająca liczba dostępnych danych eksperymentalnych to metody nie da się zastosować. W praktyce metoda jest, więc trudna do opracowania dla nowych użytkowników w krajach wdrażających swoje programy jądrowe. Jej wykorzystanie wiąże się więc najczęściej z podpisaniem umowy na wsparcie techniczne od autorów metody i wykorzystanie platformy CIAU (Code with capability of Internal Assessment of Uncertainty) [38], [39] która łączy zarówno teoretyczne podstawy metody jak również metody obliczeniowe i oprogramowanie umożliwiające jej implementację dla wybranych kodów obliczeniowych.

2.2.1. *Benchmark BEMUSE*

W efekcie rozwoju i wdrażania obliczeń z zastosowaniem metody BEPU eksperci i specjaliści pod przewodnictwem OECD NEA zaczęli organizować wspólne benchmarki (analizy porównawcze) oraz programy, aby dzielić się doświadczeniami, wiedzą oraz wspólnie badać dojrzałość metod. Program BEMUSE (**B**est **E**stimate **M**ethods – **U**ncertainty and **S**ensitivity **E**valuation) [40] pozwolił na pierwsze usystematyzowane porównanie obu dostępnych metod, pogłębioną analizę wpływu użytkownika oraz sformułowanie wniosków i rekomendacji na podstawie poczynionych obserwacji. Program podzielony był na dwie główne części, a każda z nich na trzy fazy. W pierwszej części wykonano analizy niepewności i wrażliwości dla instalacji testowej LOFT a następnie w drugiej części takie same analizy stosując te same metody dla awarii typu Large Break LOCA w elektrowni jądrowej typu PWR.

Na pierwszą część składały się trzy fazy:

1. wybór metody BEPU – tylko Uniwersytet w Pizie stosował swoją metodę propagacji niepewności wyjściowych, pozostali uczestnicy benchmarku (9 organizacji) stosowali metodę GRSu,
2. wykonanie obliczeń referencyjnych dla testu L2-5 (duże rozerwanie rurociągu zimnej nitki obiegu pierwotnego) w instalacji eksperymentalnej LOFT,
3. wykonanie obliczeń BEPU dla testu L2-5 w LOFT i oszacowanie uzyskanych niepewności.

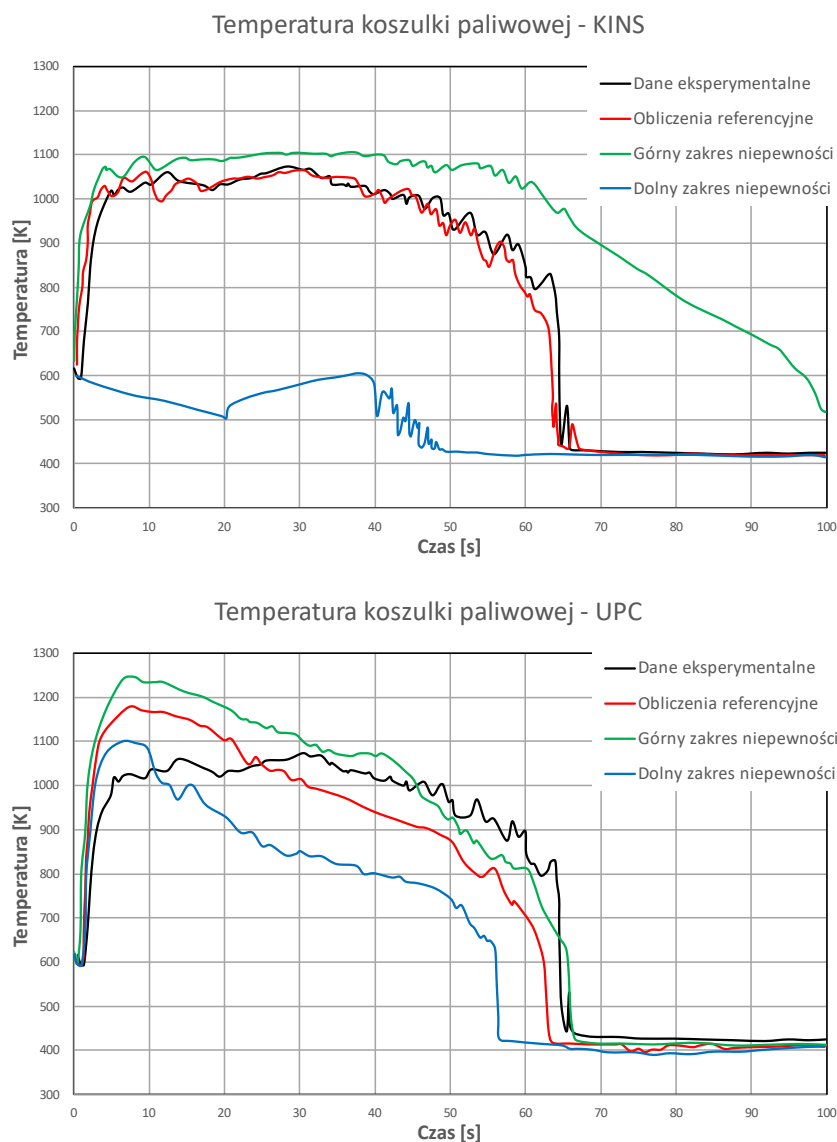
Na drugą część składały się kolejne trzy fazy:

4. wykonanie referencyjnych obliczeń najlepszego szacowania dla awarii LBLOCA w reaktorze typu PWR,
5. wykonanie obliczeń BEPU dla awarii LBLOCA w reaktorze typu PWR dla modelu elektrowni ZION (elektrownia ta znajduje się w stanie likwidacji od 1998 roku),
6. zebranie wniosków oraz rekomendacji dla użytkowników kodów i metod.

W obliczeniach BEPU dla danych z eksperymentu LOFT uczestnicy benchmarku mogli sami określić listę parametrów wejściowych które zostaną wykorzystane w analizie. Tym sposobem można było porównać wybory użytkowników, ich podejście do określenia rozkładów gęstości prawdopodobieństwa oraz określania istotnych parametrów kalibracyjnych modeli. Uczestnicy którzy wykorzystywali systematyczny proces identyfikacji oraz kategoryzacji najistotniejszych zjawisk fizycznych dla określonej awarii (ang. „Phenomena Identification and Ranking Table” – dalej PIRT[41]), określali liczbę parametrów wejściowych w zakresie od 13 do 24. Natomiast pozostali uczestnicy, którzy nie korzystali z tego procesu identyfikowali między 31 a 64 parametrów.

Na rysunku 2.3. [42] zaprezentowano porównanie uzyskanych zakresów niepewności przez dwa zespoły uczestników z KINS (Korea Południowa) oraz UPC (Hiszpania) stosujących ten sam kod obliczeniowy RELAP5. Mimo stosowania tego samego kodu oraz danych geometrycznych opisujących instalację eksperymentalną, użytkownicy uzyskali różne wyniki obliczeń referencyjnych (kolor czerwony), jak również górnego zakresu niepewności (kolor zielony) oraz dolnego zakresu niepewności (kolor niebieski). Górny zakres niepewności dla UPC po 45 sekundach jest poniżej danych doświadczalnych, co wskazuje, że wyniki te są niepoprawne, gdyż nie zapewniają odpowiedniego marginesu bezpieczeństwa. Podobnie dolny zakres niepewności dla KINS wydaje się niepoprawny, nie odzwierciedlając rzeczywistego

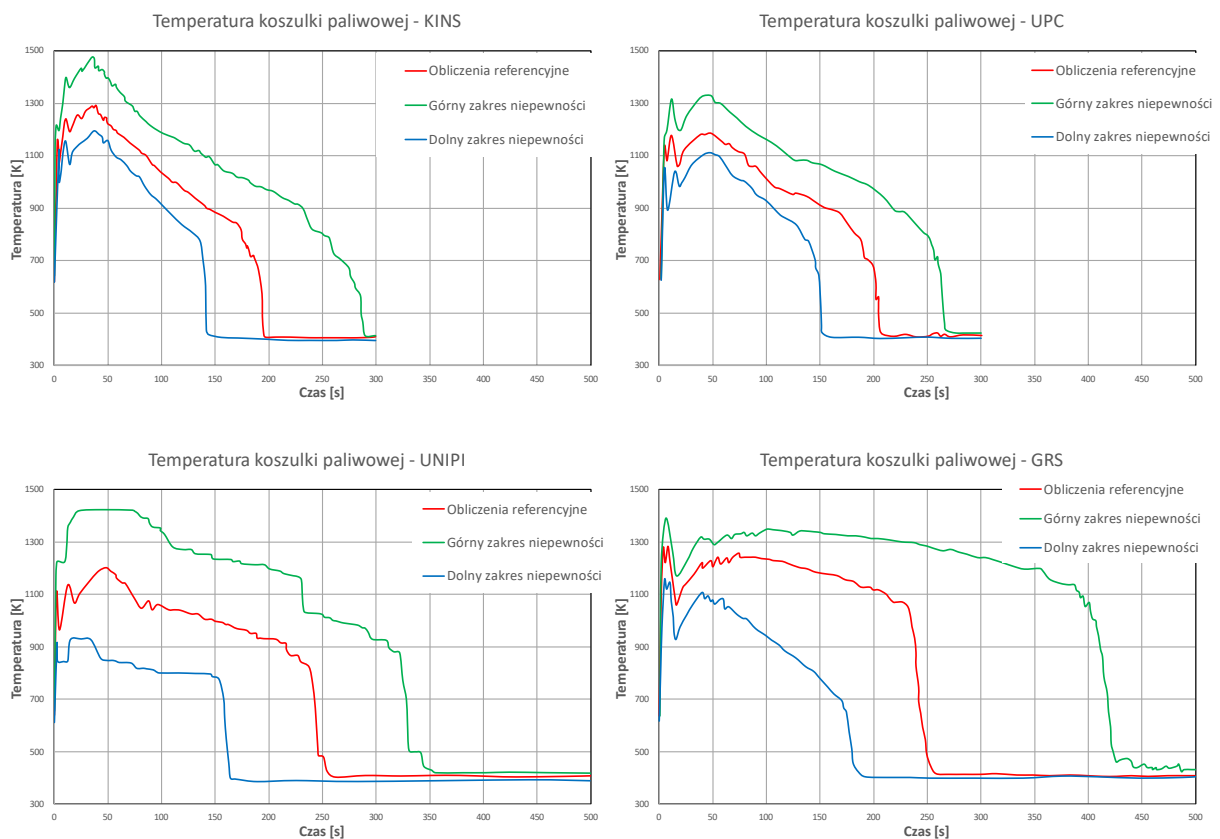
przebiegu zmian wartości wielkości wyjściowej w czasie. Rozbieżności te, podobnie jak znaczące różnice w uzyskanych wynikach przez dwie organizacje biorące udział w benchmarku wskazują na znaczący wpływ użytkownika na uzyskiwane rezultaty.



Rys. 2.3. Porównanie uzyskanych zakresów niepewności dla temperatury koszulki paliwowej dla instalacji LOFT, obliczenia wykonane przez KINS (górny rysunek) oraz UPC (dolny rysunek) [42]

W obliczeniach BEPU dla scenariusza LBLOCA w elektrowni jądrowej uczestnicy uzgodnili jedną wspólną listę parametrów wejściowych do analizy wraz z określeniem zakresów zmienności oraz rozkładów gęstości prawdopodobieństwa dla każdego z parametrów. Większość uczestników podążając za rekomendacjami z poprzedniej fazy benchmarku zwiększyła również liczbę obliczeń, zgodnie z wyższymi stopniami formuły Wilksa (tabela 2.1). Podobnie jak w poprzednim kroku wielkością badaną była temperatura koszulki paliwowej. Na rysunku 2.4. zaprezentowano porównanie uzyskanych zakresów niepewności

przez trzy zespoły uczestników z KINS (Korea Południowa), UPC (Hiszpania) oraz UNIPI (Włochy) stosujących ten sam kod obliczeniowy RELAP5 oraz GRS stosujące własny cieploprzepływowy kod ATHLET [43]. UNIPI jako jedyne stosowało metodę ekstrapolacji niepewności wyjściowych. Zastosowanie takich samych parametrów wejściowych i takich samych rozkładów przyniosło efekty, ponieważ zakresy niepewności są porównywalne dla KINS i UPC. Zakresy uzyskane przez GRS z zastosowaniem innego kodu, mogą jednak sugerować, że metoda jest mocno zależna od wykorzystywanego kodu. Niemniej jednak wciąż kluczowe jest wykonanie poprawnych obliczeń referencyjnych, ponieważ dla obliczeń KINS górny zakres niepewności jest bardzo zbliżony do limitu dla PCT wynoszącego 1477 K. Zakresy uzyskane przez UNIPI stosując inną metodę są o wiele szersze, jednak co najważniejsze wyniki mieszczą się w kryterium akceptacji. Należy podkreślić, że wykonane obliczenia porównawcze nie służyły do oceny poszczególnych zespołów analitycznych, ale do wskazania mocnych i słabych stron poszczególnych metod obliczeniowych.



Rys. 2.4. Porównanie uzyskanych zakresów niepewności dla temperatury koszulki paliwowej dla elektrowni jądrowej, obliczenia wykonane przez KINS (z lewej strony góra), UPC (z prawej strony góra) oraz UNIPI (lewa strona dół) i GRS (prawa strona dół) [44]

Wnioski dotyczące wykorzystywanych metod sformułowane na podstawie benchmarku:

- przy zwiększonej liczbie obliczeń (wyższy stopień formuły Wilksa) rozproszenie wyników jest mniejsze,
- wykonywanie większej liczby obliczeń prowadzi do bardziej wiarygodnych wyników badań wrażliwości, ponieważ zmniejsza się częstość występowania pozornych korelacji między badanymi parametrami,
- ważniejsze od liczby obliczeń są poprawność obliczeń referencyjnych oraz wybór odpowiednich parametrów wejściowych do analizy wraz z ich zakresami i rozkładami,
- obserwowano duże rozbieżności w wynikach dla obliczeń referencyjnych zarówno dla użytkowników stosujących te same kody, jak i stosujących różne kody, co jest zbliżone z wynikami wszelkich poprzednich benchmarków, efektu tego nie da się zredukować poprzez stosowanie metod BEPU,
- umiejętności, doświadczenie i wiedza użytkowników kodów na temat zarówno stosowanego kodu obliczeniowego jak i metody BEPU ma istotny wpływ na jakość wykonywanych obliczeń.

Dodatkowe rekomendacje dla użytkowników oraz nowe kierunki badań sformułowane na podstawie benchmarku:

- stosując metodę GRS należy wykorzystywać proste próbkowanie losowe (ang. simple random sampling),
- kluczowa jest kwantyfikacja niepewności parametrów modeli wykorzystywanych w kodach obliczeniowych (w tym parametrów kalibracyjnych), co najlepiej osiągnąć poprzez organizację kolejnych benchmarków bazujących na danych eksperymentalnych dostępnych dla użytkowników różnych kodów z wielu ośrodków badawczych.

2.2.2. *Benchmark PREMIUM*

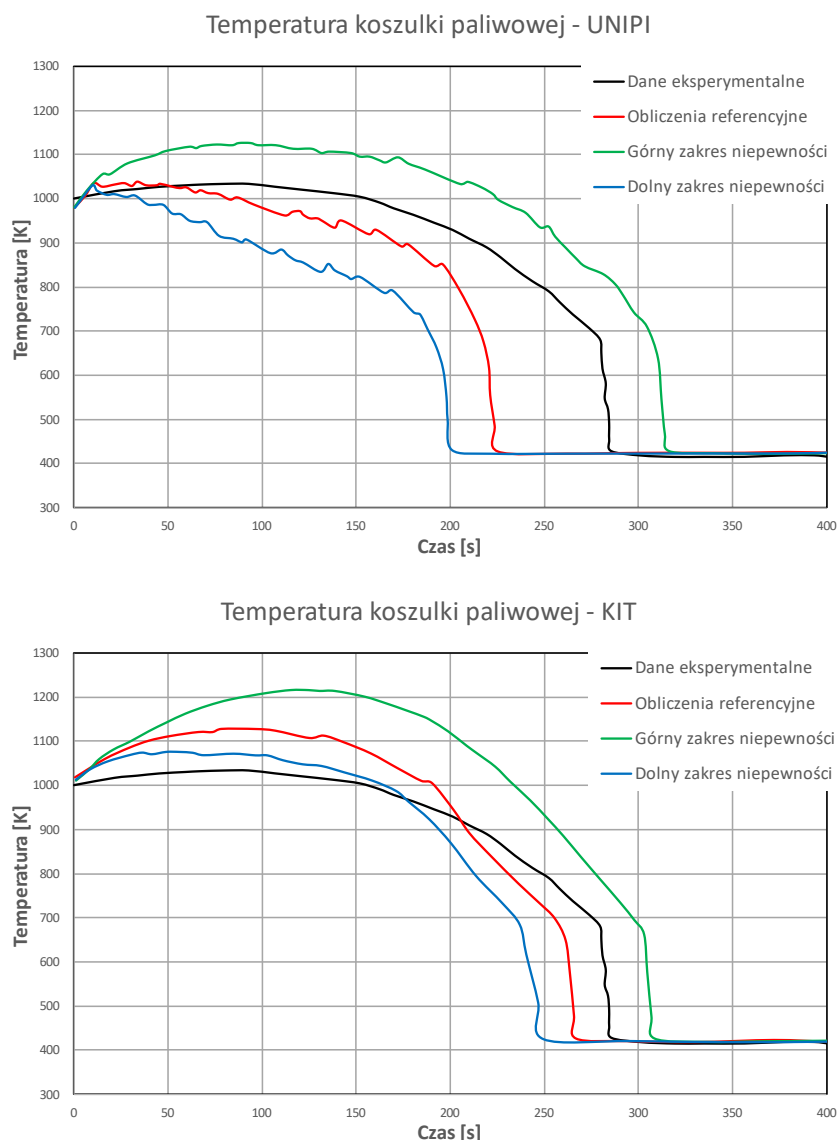
Naturalną konsekwencją wynikającą ze sformułowania rekomendacji z programu BEMUSE było rozpoczęcie kolejnego programu OECD pod nazwą PREMIUM [36], [37], [45] (**P**ost-BEMUSE **R**eflood **M**odel **I**nput **U**ncertainty **M**ethods). Program ten miał na celu przetestowanie istniejących metod kwantyfikacji niepewności parametrów modeli fizycznych stosowanych w komputerowych kodach obliczeniowych. Badane były zjawiska fizyczne zachodzące podczas fazy zalewania rdzenia (ang. „reflood”) w trakcie awarii typu LBLOCA. Uczestnicy benchmarku wykonywali obliczenia wykorzystując dane doświadczalne z dwóch

instalacji eksperymentalnych – FEBA oraz PERICLES. Obliczenia dla instalacji eksperymentalnej FEBA stanowiły punkt wyjścia do wykonania wstecznej kwantyfikacji niepewności. Na podstawie porównania stopnia zgodności wykonanych obliczeń z danymi eksperymentalnymi, szacowano wartości oczekiwane parametrów modeli wraz z ich niepewnościami.

Uczestnicy mieli do wyboru dwie metody wstecznej kwantyfikacji niepewności – metodę CIRCE opracowaną przez CEA (Francja) oraz FFTBM [35] opracowaną przez Uniwersytet w Pizie (Włochy), bądź mogli stosować własne metody (wcześniej stosowane bądź opracowane w trakcie realizacji benchmarku). Następnie uczestnicy dokonywali określenia istotnych parametrów modeli dostępnych w kodach obliczeniowych, które będą miały największy wpływ na zjawiska fizyczne podczas eksperymentów z zalewaniem rdzenia i jednocześnie starali się uzgodnić listę parametrów tak aby nie było znaczących różnic w liczbie badanych wielkości. Wykonano również obliczenia wrażliwości, aby ograniczyć liczbę parametrów do około dziesięciu. Kolejnym krokiem było wykonanie obliczeń referencyjnych dla instalacji FEBA które porównywano z danymi eksperymentalnymi. Większość uczestników uzyskała satysfakcjonującą zgodność z danymi eksperymentalnymi dla temperatury koszulki paliwowej oraz czasu, w którym dochodziło do ponownego zalania rdzenia wodą. W kolejnej fazie benchmarku uczestnicy dokonywali wstecznej kwantyfikacji niepewności. Uzyskiwane przez uczestników zakresy i rozkłady niepewności były silnie zależne od stosowanej metody, wielkości wyjściowych stanowiących podstawę porównania z eksperymentem oraz od stosowanego kodu. Po raz kolejny uwydatniony został efekt użytkownika. W ostatniej fazie benchmarku dokonano sprawdzenia uzyskanych rozkładów gęstości prawdopodobieństwa dla instalacji FEBA (testy z innymi warunkami niż w poprzedniej fazie) i instalacji PERICLES. Warto zaznaczyć, że dla instalacji PERICLES obliczenia były wykonywane bez wcześniejszej znajomości wyników eksperymentu i bez możliwości poprawy rezultatów tak aby poprawić zgodność z danymi doświadczalnymi. Jest to podejście tzw. „blind simulation”, które przypomina sytuację, w której modeluje się reaktor energetyczny i wykonuje obliczenia dla założonego scenariusza awaryjnego bez możliwości weryfikacji obliczeń z danymi pomiarowymi.

W obliczeniach sprawdzających dla instalacji FEBA większość uczestników otrzymała zakresy niepewności wielkości wyjściowej pokrywające w całości lub w części dane eksperymentalne. Na rysunku 2.5 zaprezentowano przykładowe wyniki dla jednego z uczestników który otrzymał poprawne zakresy niepewności oraz dla innego który pokrywał dane eksperymentalne tylko

w połowie czasu trwania eksperymentu. Uczestnicy ci stosowali tą samą metodę wstecznej kwantyfikacji niepewności - FFTBM.



Rys. 2.5. Zakresy niepewności dla temperatury koszulki paliwowej dla dwóch różnych uczestników stosujących tą samą metodę FFTBM w obliczeniach dla instalacji FEBA[36]

Następnie te same rozkłady parametrów wejściowych zastosowano do obliczeń weryfikacyjnych dla instalacji PERICLES. W tym wypadku uzyskane zakresy niepewności pokrywały dane eksperymentalne dla mniej niż połowy uczestników benchmarku. W przypadku rezultatów które pokrywały dane eksperymentalne uzyskiwane zakresy niepewności były szerokie lub bardzo szerokie. Uzyskiwane wyniki były niezależne od wykorzystywanego kodu obliczeniowego, a wykazywały dużą zależność od stosowanej metody. Podobnie jak w przypadku benchmarku BEMUSE zwracano uwagę na ważną rolę obliczeń referencyjnych, które w wielu wypadkach w małym stopniu pokrywały się z danymi

eksperymentalnymi. Wyniki walidacji nie były satysfakcjonujące, wskazując, że bezpośrednie zastosowanie niepewności parametrów wejściowych uzyskanych na podstawie obliczeń dla jednej instalacji, nie dają gwarancji, że dla tego samego zjawiska badanego w innej instalacji eksperymentalnej pozwolą one na uzyskanie zakresów niepewności badanych wielkości wyjściowych, które pokrywałyby dane eksperymentalne.

Spośród najczęściej stosowanych metod, FFTBM pozwalało na uzyskiwanie szerszych zakresów oraz rozkładów gęstości prawdopodobieństwa parametrów, co następnie przekładało się na szersze zakresy niepewności wielkości wyjściowych. Takie wyniki zwiększają szansę na pokrycie danych eksperymentalnych, jednak zakresy te nie powinny być zbyt szerokie, ponieważ muszą mieć podstawy fizyczne, oraz nieść wartościowe informacje predykcyjne.

Rekomendacje z projektu PREMIUM stanowiły założenia i cele kolejnego projektu OECD - SAPIUM i zostaną szczegółowo opisane w [rozdziale 3](#) przedstawiającym opracowaną metodykę, która bazuje na wnioskach z projektu SAPIUM.

3. Opis problemu badawczego – metodyka wykonywania obliczeń najlepszego szacowania wraz z oceną niepewności

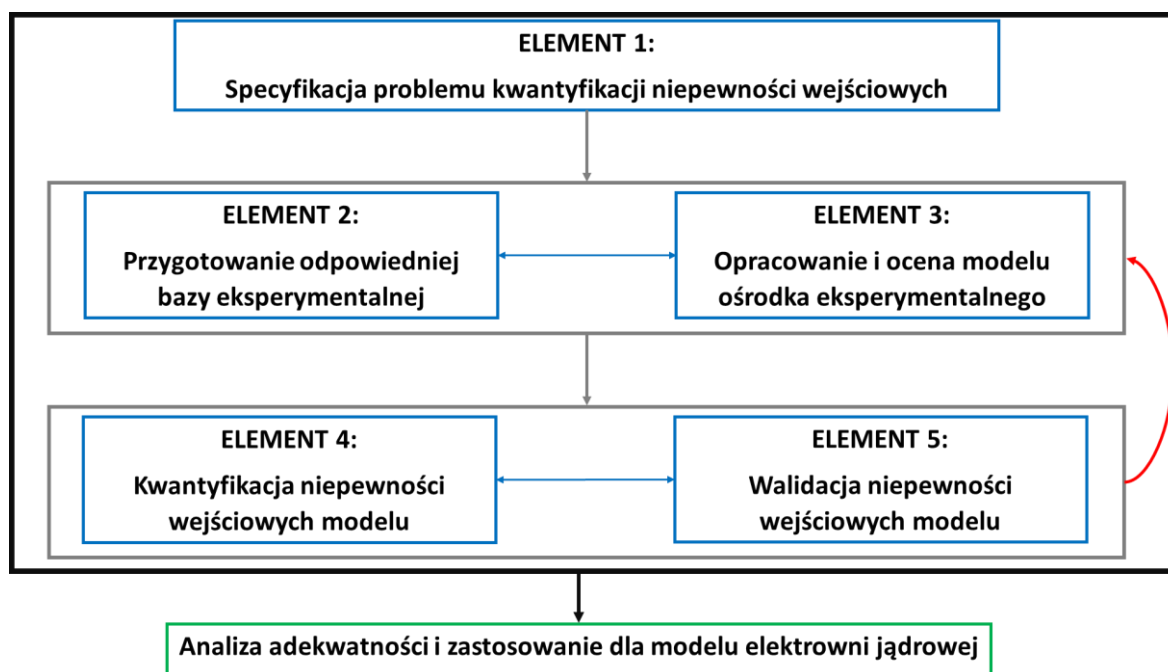
Wnioski i doświadczenia z międzynarodowych benchmarków dotyczących wykorzystania metod najlepszego szacowania wraz z oceną niepewności wskazują na konieczność ustandaryzowania procesu wykonywania takich obliczeń. Proces ten powinien obejmować między innymi kwantyfikację niepewności związanych z wewnętrznymi parametrami modeli matematycznych dostępnych w kodach obliczeniowych. We wnioskach z projektów i benchmarków zwracano uwagę, że istotny jest wybór odpowiedniej bazy eksperymentalnej do przeprowadzenia kwantyfikacji. Ma ona wpływ na zakres stosowalności uzyskanych niepewności. Nadrzędnym celem powinno być przede wszystkim ograniczenie wpływu użytkownika na uzyskiwane zakresy niepewności oraz zapewnienie ich powtarzalności. Należy również mieć na uwadze, że metody kwantyfikacji niepewności nie powinny być stosowane do warunków początkowych, brzegowych, właściwości materiałowych oraz innych parametrów fizycznych dla których możliwe jest oszacowanie ich niepewności na drodze pomiaru lub ich szacowania na podstawie statystyki.

Problemem związanym ze wsteczną kwantyfikacją niepewności jest fakt, że jest to zagadnienie, któremu brak jednoznacznej odpowiedzi. Problem ten można minimalizować poprzez upraszczanie stopnia skomplikowania zagadnienia, stąd najlepiej wykonywać kwantyfikację niepewności z wykorzystaniem prostych instalacji testowych w których badane jest jedno zjawisko. Z drugiej strony może to prowadzić do zawężania zakresu stosowalności zaproponowanych rozkładów gęstości prawdopodobieństwa parametrów modelowych kodów obliczeniowych. Stąd konieczne jest przeprowadzenie walidacji z wykorzystaniem innej instalacji, w którym badane jest to samo zjawisko. Instalacja ta może być w skali podobnej bądź w większej co jest zazwyczaj preferowane, ponieważ wskazuje czy proces skalowania może zostać następnie ekstrapolowany do zastosowania dla reaktora energetycznego. Niemniej jednak należy zaznaczyć, że dla niektórych parametrów które mają wpływ na wiele zjawisk oraz wielkości wyjściowych niemożliwe jest przeprowadzenie kwantyfikacji niepewności stosując dane z prostych instalacji testowych i wówczas należy wykonać obliczenia z wykorzystaniem bardziej złożonych instalacji ITF (Integral Test Facility).

Kolejnym problemem jest zależność niepewności jednego parametru wyjściowego od liczby analizowanych parametrów wejściowych. Stąd konieczne jest opracowanie metod preselekcji istotnych parametrów oraz poznania ich wzajemnej kowariancji. Wzrost zdolności obliczeniowych systemów komputerowych pozwala na wykorzystanie globalnych analiz

wrażliwości [46] które umożliwiają szacowanie wpływu parametrów na wielkości wyjściowe w systemach/modelach zarówno liniowych jak i nieliniowych.

Mając świadomość wszystkich powyżej opisanych problemów jak również obecnych trendów w ramach pracy doktorskiej zaczęto opracowywać metodykę pozwalającą na usprawnienie obecnych metod BEPU. Konieczne było również opracowanie metod wstecznej kwantyfikacji niepewności oraz globalnej analizy wrażliwości. W tym samym czasie publikowane były pierwsze raporty z projektu SAPIUM [16], [47] którego autorzy wskazywali, że celem projektu nie jest przygotowanie gotowej metodyki, lecz w większym stopniu zebranie obecnych trendów, dobrych praktyk, wytycznych i usystematyzowanie procesu wykonywania i walidacji obliczeń wstecznej kwantyfikacji niepewności, co powinno posłużyć odbiorcom raportu do opracowania własnych metodyk. Biorąc pod uwagę zbieżność celów i elementów metodyki oraz jej formalizację, postanowiono zaadaptować zaproponowane w ramach projektu SAPIUM elementy podejścia do BEPU na potrzeby niniejszej rozprawy doktorskiej. Wprowadzono konieczne modyfikacje oraz rozpoczęto proces opracowywania metod obliczeniowych które będą wykorzystane w metodyce a opisane w dalszej części rozprawy. Ostatnim krokiem było wdrożenie metodyki co na obecną chwilę nie zostało wykonane przez uczestników programu SAPIUM i potwierdzone odpowiednią publikacją.



Rys. 3.1. Schemat głównych elementów metodyki wykonywania obliczeń wstecznej kwantyfikacji niepewności na podstawie programu SAPIUM

Metodyka prezentuje systematyczne podejście do kwantyfikacji niepewności parametrów wejściowych modeli. Składa się z pięciu podstawowych elementów, przedstawionych na rysunku 3.1, które są następnie podzielone na bardziej szczegółowe kroki. Efektem końcowym implementacji metodyki jest możliwość wykonania obliczeń BEPU dla modelu elektrowni jądrowej, posiadając pewność, że ekstrapolacja niepewności wejściowych parametrów modeli jest możliwa. Element ten nie był częścią programu SAPIUM, stąd na rysunku 3.1. stanowi odrębny blok. W niniejszej rozprawie jest to jednak nieodłączna część opracowywanej metodyki.

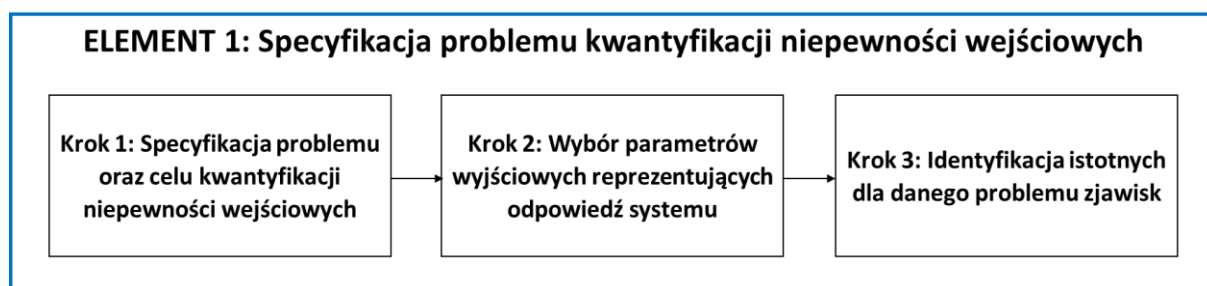
Element pierwszy jest odzwierciedleniem standardowej praktyki stosowanej w deterministycznych analizach bezpieczeństwa do określania badanego problemu. Elementy drugi oraz trzeci dostarczają podstawowych narzędzi i informacji do kwantyfikacji niepewności wejściowych. Elementy drugi oraz trzeci wzajemnie się uzupełniają i mogą być wykonywane kilkakrotnie. Jakość opracowanych w elemencie trzecim modeli oraz wykonanych obliczeń referencyjnych ma kluczowy wpływ na poprawność wykonania kolejnych elementów oraz na wysoką pewność uzyskiwanych w trakcie ich implementacji wyników. Element czwarty to wykorzystanie metod obliczeniowych do przeprowadzenia wstecznej kwantyfikacji niepewności. W elemencie piątym dokonuje się walidacji uzyskanych zakresów niepewności. W przypadku konieczności usprawnienia wyników należy wrócić do elementów 2 oraz 3 co jest oznaczone czerwoną strzałką na rysunku 3.1. Kolejne podrozdziały w szczegółowy sposób opisują kolejne kroki metodyki.

3.1. Element pierwszy metodyki – specyfikacja problemu kwantyfikacji niepewności wejściowych

Pierwszym zadaniem w każdej metodyce BEPU jest określenie problemu, czyli wybór analizowanego scenariusza awarii, określenie wielkości opisujących odpowiedzi systemów obiektu jądrowego na analizowany scenariusz oraz określenie najważniejszych zjawisk fizycznych, które będą miały kluczowy wpływ na przebieg awarii. Na element pierwszy metodyki składają się trzy podstawowe kroki przedstawione na rysunku 3.2.

W pierwszym kroku identyfikuje się typ elektrowni jądrowej która będzie analizowana (PWR, BWR, PRHR, SMR) oraz określa się scenariusz awaryjny który będzie analizowany. Istotne jest również zdefiniowanie celu wykonywania analizy. Przykładowo dla awarii typu LBLOCA można prowadzić analizę weryfikującą odpowiedź systemów bezpieczeństwa lub analizę weryfikującą utrzymanie wytrzymałości materiałowej obudowy bezpieczeństwa. Wybór scenariusza oraz celu wykonywania analizy ma wpływ na kryteria akceptacji oraz analizowane

parametry najlepiej reprezentujące odpowiedź systemu. W przypadku wielu różnych parametrów wyjściowych konieczne jest przeprowadzenie kwantyfikacji niepewności dla większej liczby parametrów wejściowych, co znacznie zwiększa nakład prac, czas oraz koszty obliczeniowe. Sugerowane jest więc grupowanie scenariuszy tak aby wewnątrz grupy znajdowały się różne scenariusze, dla których badane będą te same wielkości wyjściowe. Przykładem takiej grupy są wszelkie scenariusze awaryjne w których dochodzi do zmniejszenia ilości wody w obiegu chłodzenia reaktora, tj. awarie typu LOCA czy rozerwanie rurek wytwornicy pary.



Rys. 3.2. Kroki należące do pierwszego elementu metodyki

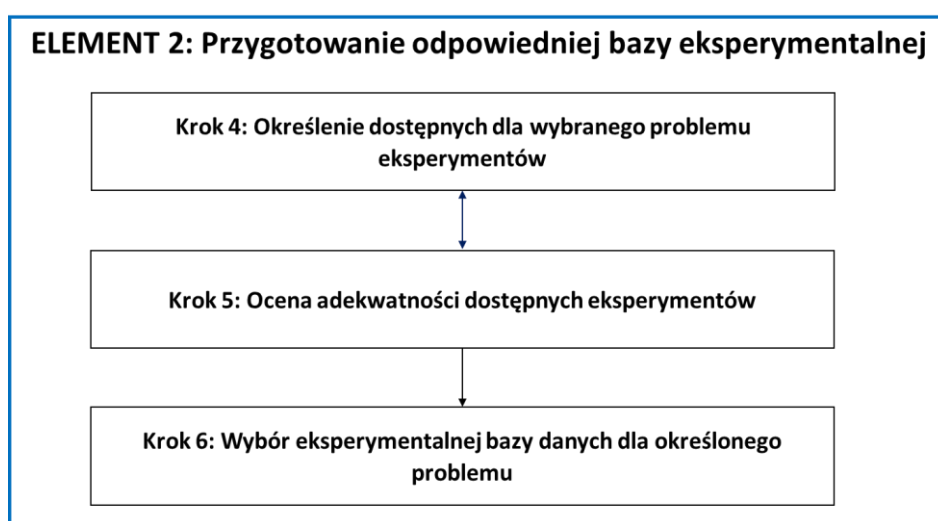
Drugi krok metodyki polega na określeniu wielkości wyjściowych reprezentujących odpowiedź systemu na określony w kroku pierwszym scenariusz awaryjny. Wykonując deterministyczne analizy bezpieczeństwa bada się parametry które pozwalają na weryfikację kryteriów akceptacji, takie jak temperatura koszulki paliwowej, wartość parametru DNBR (ang. „departure from nucleate boiling ratio”), stopień utlenienia koszulek paliwowych czy ciśnienie w obiegu pierwotnym. Często stosuje się również parametry zastępcze, nie odnoszące się wprost do kryteriów akceptacji, ale pozwalające wnieść informacje, które pozwolą na oszacowanie wypełnienia kryteriów akceptacji. Takimi parametrami są ilość chłodziwa w obiegu czy poziom wody w zbiorniku ciśnieniowym reaktora. W stosowanej metodyce parametry wyjściowe reprezentujące odpowiedź systemu muszą być dostępne bezpośrednio i mierzalne w eksperymentach określonych w elemencie drugim metodyki.

W trzecim kroku metodyki dokonuje się wyboru najistotniejszych zjawisk fizycznych które będą miały wpływ na przebieg scenariusza awaryjnego. Nie wszystkie zjawiska mają taki sam wpływ na przebieg awarii oraz jej skutki, stąd ich ograniczenie jest uzasadnione. Jeśli analizowany scenariusz posiada znane fazy rozwoju warunków awaryjnych, wówczas należy przeprowadzić proces ograniczenia liczby zjawisk i związanych z nimi parametrów dla każdej z faz. Należy też podzielić zjawiska w zależności od komponentów, dla których zjawiska te są istotne. Wykonuje się to poprzez wykonanie procesu zwieńczonego opracowaniem tabeli

szeregującej zjawiska - PIRT[41]. W pracach [23], [48] zidentyfikowano 116 istotnych zjawisk ciepłno-przepływowych występujących podczas typowych scenariuszy awaryjnych. Dokładniejszy opis procesu PIRT znajduje się w [rozdziale 4.1](#) gdzie proces ten jest zaprezentowany w praktycznym zastosowaniu.

3.2. Element drugi metodyki – przygotowanie odpowiedniej bazy eksperymentalnej

Celem kolejnego elementu metodyki jest przygotowanie reprezentatywnej bazy eksperymentalnej dla problemu zdefiniowanego w elemencie pierwszym. Proces ten wykonuje się w trzech krokach zaprezentowanych na rysunku 3.3.



Rys. 3.3. Kroki należące do drugiego elementu metodyki

Krok czwarty metodyki to określenie listy instalacji eksperymentów dostępnych dla określonego uprzednio problemu/scenariusza awaryjnego. Krok ten jest warunkowany możliwością dostępu do baz danych eksperymentalnych. W krajach które jako pierwsze wdrażały energetykę jądrową, w latach osiemdziesiątych prowadzone były intensywne kampanie eksperymentalne, stąd posiadają one obszerne bazy danych eksperymentalnych. Jedynie część z tych danych jest ogólnodostępna lub dostępna w bazie danych OECD NEA, która jest najczęściej wykorzystywanym źródłem danych eksperymentalnych. Zebrane są tam zarówno dane z instalacji SETF jak i ITF. Należy więc dążyć, aby dla określonego scenariusza awaryjnego oraz jego faz odnaleźć odpowiadające eksperymenty, w których badane były najważniejsze zjawiska określone w kroku trzecim. Jeśli nie jest to możliwe, można założyć, że zjawisko badane w instalacji eksperymentalnej dla jednej z faz eksperymentu będzie można wraz z towarzyszącymi mu niepewnościami przełożyć na pozostałe fazy. Dla wybranych eksperymentów powinno się opisać jakie modele wykorzystywane w kodzie obliczeniowym

mogą zostać zweryfikowane przy zastosowaniu danych eksperymentalnych. Na potrzeby procesu skalowania należy również opisać warunki fizyczne i geometryczne.

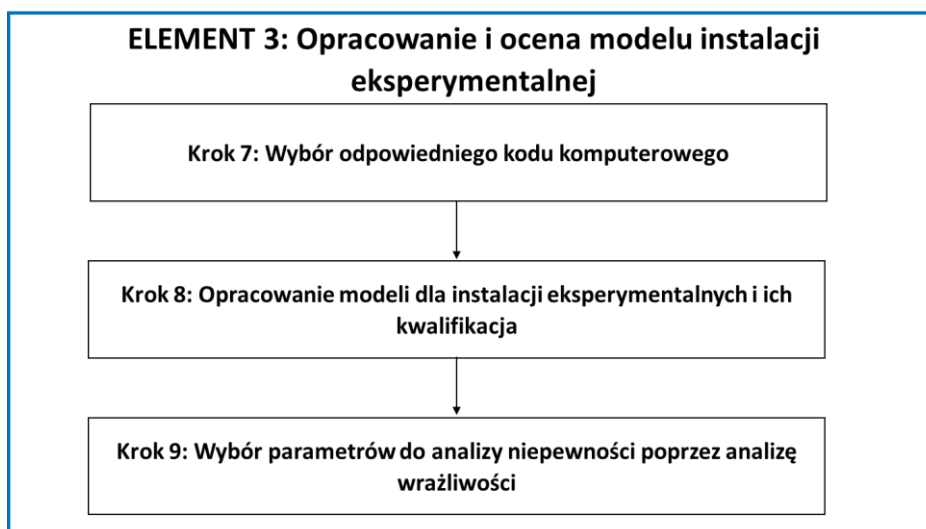
W kroku piątym weryfikuje się adekwatność bazy danych pod względem dwóch właściwości: reprezentatywności oraz kompletności. Reprezentatywność oznacza, że wybrane eksperymenty dostarczają istotnych informacji na podstawie których będzie można dokonać kwantyfikacji i walidacji niepewności parametrów wejściowych. Kompletność zapewnia, że całe spektrum zjawisk definiujących scenariusz awaryjny jest pokryte przez zidentyfikowaną w kroku czwartym bazę eksperymentalną. W przeprowadzeniu obu tych procesów pomocne są matryce zjawisk dla instalacji SETF [15] oraz ITF [14] opracowane przez ekspertów pod egidą OECD NEA. Reprezentatywność ocenia się względem: jakości i dostępności danych eksperymentalnych, odpowiedniej reprezentacji zjawiska przez mierzone dane eksperymentalne, informacji dotyczących niepewności pomiarowych wielkości wyjściowych, warunków początkowych i brzegowych w odniesieniu do warunków w problemie zidentyfikowanym w kroku pierwszym. Kompletność w największym uproszczeniu to taka właściwość przygotowanej bazy danych eksperymentalnych, gdy dodawanie kolejnych danych nie zwiększa poziomu jej adekwatności dla zdefiniowanego problemu.

Na podstawie procesu PIRT (krok 3) oraz analizy adekwatności (krok 5) dokonuje się wyboru końcowej bazy danych eksperymentalnych zarówno na potrzeby kwantyfikacji niepewności parametrów wejściowych jak również ich późniejszej walidacji. Jest to opisane przez krok szósty metodyki. Zalecane jest, aby stosować instalacje eksperymentalne w mniejszej skali (SETF) dla kwantyfikacji niepewności wejściowych wybranych modeli kodów obliczeniowych. Należy dążyć, aby stosować co najmniej dwa eksperymenty dla jednego zjawiska i modelu kodu obliczeniowego, aby zapewnić zwielokrotnienie i różnorodność danych eksperymentalnych. Na potrzeby walidacji zalecane jest wykorzystywanie instalacji eksperymentalnych w większej skali w których badane było kilka zjawisk w tym samym czasie, czyli instalacji typu ITF. W przypadku gdy ograniczenia bazy danych eksperymentalnych nie pozwalają na sugerowany podział bazy, należy zastosować wszystkie dostępne eksperymenty dla kwantyfikacji niepewności, a walidację przeprowadzać dla testów z innymi warunkami początkowymi przeprowadzonymi w tych samych instalacjach.

3.3. Element trzeci metodyki – opracowanie i ocena modelu instalacji eksperymentalnej

Element trzeci składa się z trzech powiązanych ze sobą kroków, które wymagają opracowania modeli matematycznych instalacji eksperymentalnych w wybranym kodzie obliczeniowym.

Dla przygotowanego modelu instalacji wykonuje się obliczenia referencyjne i dokonuje kwalifikacji modelu na podstawie zgodności wyników obliczeń z danymi eksperymentalnymi. W ostatnim kroku dla zakwalifikowanego modelu wykonuje się analizę wrażliwości w celu selekcji najważniejszych parametrów wejściowych, których niepewności będą kwantyfikowane. Schemat elementu trzeciego przedstawiony jest na rysunku 3.4.



Rys. 3.4. Kroki należące do trzeciego elementu metodyki

W kroku siódmym należy wybrać dopasowany do potrzeb kod obliczeniowy. Dopasowanie ocenia się na podstawie przeznaczenia i zakresu stosowalności kodu oraz podstawowych modeli i korelacji stosowanych w kodzie. W przypadku obliczeń dla awarii projektowych (takich jak LOCA, rozszczelnienie rurociągu parowego czy rozerwanie rurek wytwornicy pary) stosuje się zwyczajowo systemowe kody do obliczeń ciepłno-przepływowych takie jak RELAP5[49], TRACE[50], [51] czy CATHARE[52]. Wszystkie te kody bazują na sześciu równaniach bilansu dla cieczy i pary wodnej. Równania bilansu są sprzężone z równaniami przewodzenia ciepła oraz kinetyki neutronów. Równania te uzupełniane są w każdym kodzie odmiennymi równaniami domykającymi, w zależności od zakresu ich stosowalności oraz docelowego przeznaczenia. Każdy z kodów obliczeniowych posiada więc inny zestaw wejściowych parametrów określających dedykowane, dodatkowe modele bazujące na empirycznych korelacjach.

Ósmy krok jest jednym z najbardziej istotnych w całej metodyce. W kroku tym opracowuje się modele zidentyfikowanych w elemencie drugim instalacji eksperymentalnych stosując wybrany w poprzednim kroku kod obliczeniowy. Przygotowanie modeli powinno być poprzedzone przygotowaniem ich nodalizacji wedle jasno sprecyzowanych wytycznych, tak aby nodalizacje dla modeli wszystkich instalacji były jak najbardziej do siebie zbliżone pod

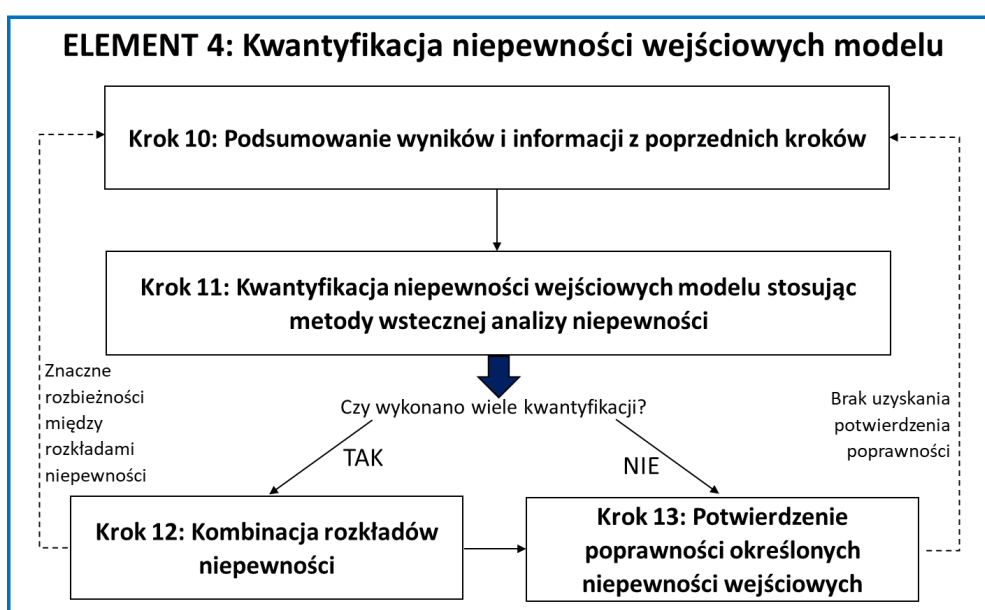
względem schematu nodów. Podobnie stosowane opcje kodu dotyczące modeli fizycznych powinny być w miarę możliwości używane jednolicie w każdym z opracowywanych modeli instalacji eksperymentalnych. Po opracowaniu modeli instalacji eksperymentalnych należy dokonać ich kwalifikacji na podstawie porównania z danymi pomiarowymi. Proces kwalifikacji dzieli się na cztery główne podprocesy. Pierwszym punktem jest demonstracja zgodności geometrycznej opracowanej nodalizacji z rzeczywistymi wymiarami. Drugim podprocesem jest porównanie wyników stanu ustalonego z warunkami początkowymi i brzegowymi określonymi i mierzonymi w eksperymencie. Weryfikacja poprawności stanu ustalonego jest niezwykle istotna, ponieważ daje pewność, że wartości parametrów są stabilne i wszelkie zmiany tych parametrów które będą obserwowane podczas stanów nieustalonych wynikają z symulacji zachowania systemów obiektu a nie z fluktuacji numerycznych powiązanych z niewłaściwym modelowaniem. Trzecim podprocesem jest jakościowa ocena rezultatów stanu nieustalonego. Dokonuje się tego poprzez porównanie wyników obliczeń z danymi eksperymentalnymi dla przebiegów czasowych wybranych wielkości wyjściowych. Należy dokonać również sprawdzenia czy sekwencja zdarzeń opisująca rozwój scenariusza awaryjnego jest zgodna z danymi eksperymentalnymi oraz czy wszelkie istotne zjawiska cieplno-przepływowe zostały zasymulowane. Dokonuje się tego najczęściej przez porównanie skalnych wartości niektórych parametrów oraz sprawdzenie czasów wyzwolenia niektórych sygnałów, np. sygnału wyzwolenia wtrysku wody z systemów awaryjnego zalewania rdzenia. Ostatnim podpunktem jest ilościowa analiza zgodności wyników obliczeń z danymi eksperymentalnymi. Można tego dokonać chociażby przez klasyczne obliczenia odchylenia standardowego czy błędu średniokwadratowego w całym zakresie czasowym bądź porównania wartości parametrów które są najistotniejsze w opisie zjawiska. Proces kwalifikacji powinien być iteracyjny, należy również pamiętać o zakresie stosowalności przygotowanego modelu.

Krokiem dziewiątym metodyki jest wybór najważniejszych parametrów wejściowych które będą następnie poddane kwantyfikacji. Proces wyboru tych parametrów wspomagany jest poprzez wykonanie analizy wrażliwości. Analizy wrażliwości najlepiej wykonać w dwóch odmiennych podejściach jedno po drugim. Wpierw należy wykonać lokalną analizę wrażliwości dla dużej liczby parametrów, aby wstępnie ograniczyć ich liczbę. Następnie należy wykonać globalną analizę wrażliwości pozwalającą na ranking istotności wpływu parametrów na wielkości wyjściowe modelu oraz ocenę korelacji między parametrami co umożliwi odrzucenie parametrów które mają wpływ na wyniki tylko poprzez skorelowanie z innym

bardziej istotnym parametrem. Zaproponowana metoda wyboru parametrów wejściowych opisana została w [rozdziale 4.4](#) niniejszej rozprawy.

3.4. Element czwarty metodyki – kwantyfikacja niepewności wejściowych modelu

W czwartym elemencie wykorzystuje się modele instalacji eksperymentalnych, opracowane w ramach elementu trzeciego, w celu przeprowadzenia kwantyfikacji niepewności wejściowych parametrów modelowych. Kroki tego elementu są zbliżone do opisanych w poprzednim rozdziale faz projektu PREMIUM. Schemat kroków składających się na element czwarty metodyki przedstawiono na rysunku 3.5.



Rys. 3.5. Kroki należące do czwartego elementu metodyki

Krok dziesiąty jest krokiem pomocniczym, gdzie zbiera się i porządkuje wszystkie dostępne z poprzednich kroków informacje, sprawdza ich poprawność oraz dokonuje ewentualnych korekt. Jedną z korekt jaką można wprowadzić jest przegląd istotności wielkości wyjściowych analizowanych we wszystkich instalacjach eksperymentalnych. Na podstawie takiego przeglądu można im przypisać różne wagi w obliczeniach wykonywanych w kolejnym kroku.

Krok jedenasty wymaga opracowania metody kwantyfikacji niepewności wejściowych parametrów modeli stosowanych w kodzie obliczeniowym. Metody te pozwalają na uzyskanie rozwiązania problemu wstecznej propagacji niepewności. W metodach tych dąży się do uzyskania zakresów i rozkładów gęstości parametrów wejściowych na podstawie określenia stopnia zbieżności między wynikami obliczeń a danymi eksperymentalnymi. W [rozdziale 4.5](#) rozprawy opisano dwie zaproponowane metody wykonywania obliczeń wstecznej niepewności

pozwalających na dokonanie procesu kwantyfikacji niepewności parametrów wejściowych, stąd tamże znajdzie się szczegółowy opis tego procesu. Jest to jeden z głównych kroków całej metodyki, prowadzący do uzyskania rozkładów gęstości prawdopodobieństwa badanych parametrów wejściowych modelu.

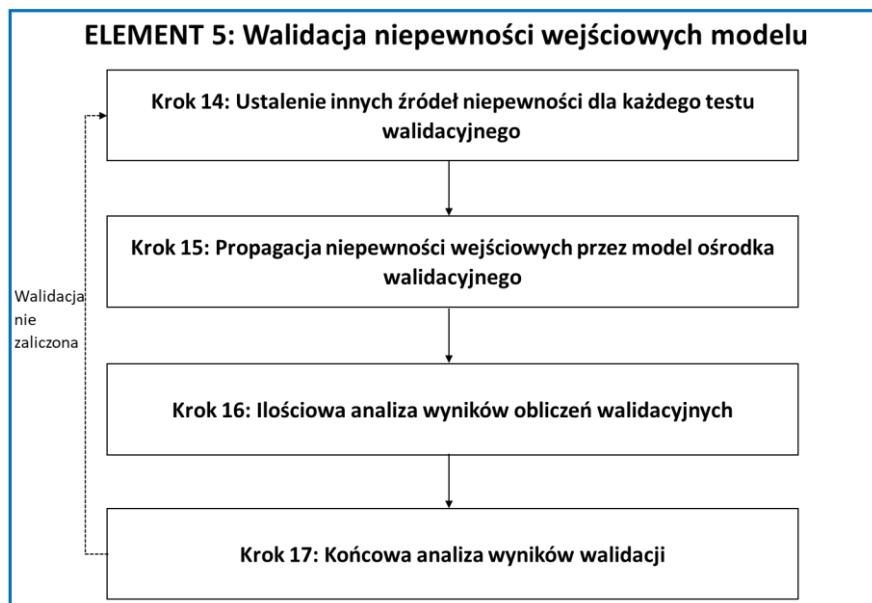
Krok dwunasty jest opcjonalny i występuje jedynie wówczas, gdy dla wybranych parametrów wejściowych wykonywane są niezależne obliczenia kwantyfikacji niepewności z zastosowaniem różnych instalacji eksperymentalnych. W takim wypadku należy więc dokonać syntezy rozkładów uzyskanych z kilku źródeł niepewności parametrów wejściowych. Najczęściej wykorzystuje się metody średnich ważonych, aby uzyskać jeden rozkład niepewności z kilku rozkładów uzyskanych dla różnych obliczeń. W przypadku gdy różnice w uzyskanych zakresach lub rozkładach niepewności nie są pomijalne, należy wrócić do kroku dziesiątego, aby dokonać analizy czy należy wrócić do poprzednich kroków czy wykonać nowe obliczenia np. z innymi założeniami.

Krok trzynasty to przeprowadzenie obliczeń sprawdzających wykorzystując uzyskane zakresy oraz rozkłady gęstości prawdopodobieństwa parametrów wejściowych. Obliczenia te wykonuje się dla testów niezastosowanych dla danej instalacji eksperymentalnej w kroku jedenastym. Testy te mogą posiadać nieznacznie zmienione warunki początkowe, brzegowe czy geometryczne, ale należą do jednej grupy zakresu stosowalności. Wykonuje się więc klasyczne obliczenia propagacji niepewności wejściowych, a uzyskane rezultaty porównuje z danymi eksperymentalnymi z uwzględnieniem oszacowanego błędu pomiarowego. Jest to krok podobny do walidacji wykonywanej w elemencie piątym, stąd może być z nim czasem łączony. Metody opisane do sprawdzenia jakości wyników w walidacji są również stosowane w kroku trzynastym.

3.5. Element piąty metodyki – walidacji niepewności wejściowych modelu

Ostatnim elementem metodyki jest walidacja wyznaczonych w elemencie czwartym niepewności wejściowych modeli stosowanych w kodzie obliczeniowym. Do walidacji wykorzystuje się modele instalacji eksperymentalnych, które nie były stosowane w elemencie czwartym, bądź testy o innych warunkach początkowych i brzegowych przeprowadzone w niektórych instalacjach wykorzystywanych w czwartym elemencie. Walidacja nie może być przeprowadzona w domenie parametrów wejściowych, ponieważ porównanie z danymi eksperymentalnymi jest wówczas niemożliwe. Stosuje się więc, obecnie wykorzystywaną metodę propagacji niepewności wejściowych przez kod obliczeniowy. W związku z tym dla

tego elementu nie było konieczności opracowywania nowych metod obliczeniowych. Na element piąty składają się cztery główne kroki przedstawione na rysunku 3.6.



Rys. 3.6. Kroki należące do piątego elementu metodyki

Jak wspomniano walidacja nie może być przeprowadzona w domenie niepewności wejściowych, lecz w odniesieniu do mierzalnych wielkości wyjściowych. Stąd w kroku czternastym należy zidentyfikować pozostałe źródła niepewności, które mogą wpływać na wyniki wykonywanych obliczeń. Dotyczy to przede wszystkim niepewności związanych z warunkami początkowymi, brzegowymi oraz niepewności modelowania których nie da się zminimalizować. Wszelkie źródła niepewności zostały opisane w [rozdziale 1.2.](#)

W kroku piętnastym wykorzystując modele instalacji eksperymentalnych przeznaczonych na cele walidacji i przygotowanych w kroku ósmym, wykonuje się obliczenia niepewności wyjściowych z zastosowaniem wyznaczonych w elemencie czwartym niepewności wejściowych parametrów modeli. Obliczenia te wykonuje się stosując metody propagacji niepewności wejściowych przez kod obliczeniowy, które zostały opisane w [rozdziale 2.1.1.](#)

W kroku szesnastym dokonywana jest ilościowa analiza wyników obliczeń walidacyjnych. Najprostsza analiza może wskazać czy dane eksperymentalne są pokryte przez uzyskany zakres niepewności badanej wielkości wyjściowej. Taka analiza premiuje jednak bardzo szerokie zakresy niepewności, które nie niosą wówczas informacji o rzeczywistym rozkładzie zmiennej wyjściowej. Konieczne jest więc zdefiniowanie dodatkowych indyktorów walidacji pozwalających na bardziej szczegółową analizę. W projekcie SAPIUM proponowano głównie

indykatory dla wielkości skalarnych, natomiast w ramach niniejszej rozprawy zaproponowane zostały w ramach metody opartej na uczeniu maszynowym opisanej w [rozdziale 4.5.2.](#)

Krok siedemnasty to zebranie wyników walidacji dla wszystkich eksperymentów i jeśli to konieczne porównanie i agregacja kilku wartości indykatorów. Agregacja podobnie jak w kroku dwunastym dokonywana jest z zastosowaniem prostej metody średnich arytmetycznych bądź ważonych. W przypadku uzyskania poprawnych wyników walidacji oraz stwierdzenia jej adekwatności można uznać, że kwantyfikacja niepewności wejściowych została zakończona pozytywnie i wyznaczone rozkłady gęstości prawdopodobieństwa mogą zostać wykorzystane do analizy w większych obiektach energetyki jądrowej.

Aby omówić najważniejsze zagadnienia związane z adekwatnością oraz ekstrapolacją wyników do zastosowań w większej skali dodano element szósty, który nie znajdował się w opisie metodyki przedstawionym w programie SAPIUM.

3.6. Element szósty metodyki – analiza adekwatności oraz ekstrapolacja do obiektu energetyki jądrowej

Problem skalowania pojawia się w zagadnieniach związanych z wykonywaniem obliczeń numerycznych z wykorzystaniem kodów obliczeniowych od samego początku ich stosowania, gdyż wiele równań wykorzystywanych w modelach matematycznych wynika z doświadczeń przeprowadzonych w mniejszej skali. Problem ten jest istotny zarówno przy opracowywaniu kodu – przy opracowywaniu modeli specjalnych dla pojedynczych zjawisk opisanych przez równania konstytutywne – jak również przy walidacji kodu lub kwantyfikacji niepewności parametrów wejściowych. Rozwiązanie tego problemu zdecydowanie wykracza poza zakres niniejszej rozprawy, wyczerpujący opis został przedstawiony w raporcie [53]. Najprostszą metodą na minimalizację wpływu skalowania i zapewnienie adekwatności uzyskanych wyników wstecznej kwantyfikacji niepewności jest wykonanie walidacji omówionej w elemencie piątym wykorzystując instalację w większej skali geometrycznej, w której badano wiele zjawisk i która posiadała warunki początkowe jak najbliżej zbliżone do tych obserwowanych w obiektach energetyki jądrowej. Problem właściwej ekstrapolacji do skali obiektów energetyki jądrowej został zidentyfikowany w raporcie SAPIUM jako jedno z kluczowych zagadnień do rozwiązania w przyszłości, między innymi poprzez organizację kolejnych benchmarków oraz projektów międzynarodowych.

Zastosowanie metodyki kończy się, więc poprzez identyfikację istotnych parametrów do analizy BEPU dla wybranego scenariusza dla obiektu energetyki jądrowej, a następnie wykonanie analizy stosując metodę propagacji niepewności wejściowych przez kod. Wartości

oczekiwane i rozkłady gęstości prawdopodobieństwa dla parametrów specjalnych modeli stosowanych w kodzie określone są na podstawie przeprowadzonej zgodnie z opisaną metodyką kwantyfikacją, pozostałe niepewności parametrów reprezentujących warunki początkowe i brzegowe są opisane przez dostępne zakresy błędów pomiarowych tych wielkości. Ocena poprawności uzyskanych zakresów niepewności wielkości wyjściowych odbywa się poprzez porównanie maksymalnych wartości z tych zakresów z kryteriami akceptacji. Krok ten jest więc zbliżony do kroków elementu piątego metodyki, z tą różnicą, że wyniki są ostatecznie porównywalne z istotnymi kryteriami akceptacji.

4. Wdrożenie pierwszych elementów metodyki wykonywania obliczeń najlepszego szacowania wraz z oceną niepewności

Główna część rozprawy poświęcona jest opracowaniu metod obliczeniowych które wykorzystywane będą do kwantyfikacji niepewności wejściowych w ramach stosowania metodyki wykonywania obliczeń najlepszego szacowania wraz z oceną niepewności. Metody te jednak muszą być sprawdzone w praktycznych zastosowaniach oraz następnie zweryfikowane. Podjęto więc próbę wdrożenia całej metodyki opisanej w [rozdziale trzecim](#). Pełne wdrożenie takiej metodyki stanowi wysiłek, który powinien podjąć cały zespół analityków mając gotowe metody obliczeniowe. Stąd w rozprawie wdrożenie metodyki odbywa się w sposób ograniczony do kilku instalacji eksperymentalnych oraz jedynie podstawowych zjawisk dla wybranej awarii, bez pokrycia całego spektrum zjawisk co byłoby oczekiwane przy pełnym wdrażaniu metodyki. Pozwoli to, jednakże na wystarczającą demonstrację użyteczności metod obliczeniowych oraz zastosowania metodyki. Zakres jest, więc dostosowany do ogólnego celu rozprawy, aby wciąż można było dokonać ogólnej oceny opracowanych metod obliczeniowych na podstawie praktycznego zastosowania. Dokonano również walidacji rozkładów gęstości prawdopodobieństwa, uzyskanych w wyniku wstecznej kwantyfikacji niepewności. W związku z tym konieczne było opracowanie modelu jednej instalacji w większej skali. W instalacji badano zjawiska dla których dokonywano wstecznej kwantyfikacji niepewności. Mając na względzie opisane ograniczenia i zakres rozprawy kolejne podrozdziały prezentują wdrożenie trzech pierwszych elementów metodyki opisanej w [rozdziale trzecim](#) oraz metody obliczeniowe które będą wykorzystywane w elementach trzecim oraz czwartym metodyki.

4.1. Wdrożenie pierwszego elementu metodyki

Finalna analiza najlepszego szacowania wraz z oceną niepewności została wykonana dla elektrowni jądrowej typu PWR, analizowanym scenariuszem było duże rozerwanie jednego z rurociągów zimnej nitki pierwotnego obiegu chłodzenia reaktora jądrowego (LBLOCA). Celem analizy jest weryfikacja czy odpowiedź systemów bezpieczeństwa jest wystarczająca do opanowania skutków awarii. Zdecydowano się na wybór takiego scenariusza awaryjnego ze względu na jasność w określeniu kryteriów akceptacji, dostępność literaturową obszernych opisów zjawisk które zachodzą podczas takiej awarii oraz fakt, że jest to awaria graniczna dla awarii projektowych dla której w wielu krajach pozwala się na wykonywanie obliczeń stosując metody BEPU[54], [55].

Głównymi wielkościami wyjściowymi które reprezentują odpowiedź systemu na założony scenariusz awaryjny dla awarii typu LBLOCA są maksymalna temperatura koszulki, maksymalny lokalny stopień utlenienia koszulek, oraz zapewnienie długoterminowego chłodzenia rdzenia, aby odbierane było ciepło powyłączeniowe, co można określać poprzez poziom wody w rdzeniu. Istotnymi parametrami pozwalającymi na ocenę rozwoju przebiegu awarii typu LBLOCA są również wypływ masowy przez rozerwanie, ciśnienie w obiegu pierwotnym oraz ilość chłodziwa w obiegu pierwotnym. Spośród wymienionych wielkości wyjściowych łatwo mierzalne w instalacjach testowych są w szczególności temperatura koszulki paliwowej, masowe natężenie wypływu, ciśnienie w obiegu pierwotnym oraz poziom wody w rdzeniu.

Dla wybranego scenariusza awarii w kolejnym kroku metodyki dokonuje się wyboru najważniejszych zjawisk fizycznych które będą miały wpływ na przebieg scenariusza awaryjnego. Proces identyfikacji najważniejszych zjawisk dla awarii typu LBLOCA jest dobrze udokumentowany literaturowo [14], [23], [48], jednak sam proces określenia istotności zjawisk jest uniwersalny stąd jego najważniejsze punkty zostaną pokrótce omówione. Rozpoczyna się on od identyfikacji wszystkich zjawisk fizycznych oraz procesów, które mogą zajść w trakcie rozpatrywanej awarii. Następnie należy określić wszystkie systemy lub elementy wyposażenia obiektu jądrowego w których zjawiska albo procesy będą obserwowalne, bądź będą miały wpływ na działanie takiego systemu bądź elementu. Kolejnym krokiem jest podział scenariusza awaryjnego na charakterystyczne przedziały czasowe, które odpowiadają fazom rozwoju awarii i w trakcie których wybrane zjawiska są obserwowane w sposób ciągły. W ten sposób można dokonać rozdziału zjawisk względem czasu ich trwania i miejsca występowania. W przypadku awarii typu LBLOCA wydzieli się trzy główne fazy awarii: wydmuchu (ang. „blowdown”), napelnienia (ang. „refill”) oraz zalewania rdzenia (ang. „reflood”). Następnie dla najistotniejszych komponentów i systemów takich jak: pręty paliwowe, rdzeń reaktora, górna komora mieszania, szczelina opadowa zbiornika reaktora, stabilizator ciśnienia, rurociągi i pompy obiegu pierwotnego, wytwornica pary, systemy bezpieczeństwa czy miejsce wystąpienia rozerwania identyfikuje się procesy i zjawiska, które mogą zajść w trakcie każdej z trzech faz awarii. W przypadku awarii typu LBLOCA bardzo wiele zjawisk zachodzi w rdzeniu reaktora oraz w szczelinie opadowej zbiornika ciśnieniowego. Związane jest to z szybką oraz dużą ucieczką chłodziwa przez rozerwanie a następnie wtryskiem wody z systemów awaryjnego zalewania rdzenia, co powoduje gwałtowne zmiany w procesach odbioru ciepła z rdzenia reaktora.

Po identyfikacji wszystkich zjawisk i procesów zachodzących w różnych fazach awarii oraz komponentów i systemów, dla których te procesy i zjawiska są obserwowane dokonuje się przeglądu oraz rankingu ich istotności w odniesieniu do określonych wcześniej istotnych wielkości wyjściowych analizy. W poprzednim kroku zidentyfikowano cztery najważniejsze wielkości wyjściowe takie jak temperatura koszulki paliwowej, poziom wody w rdzeniu, masowe natężenie wypływu oraz ciśnienie w obiegu pierwotnym, stąd ocena istotności zjawisk fizycznych i procesów odnosić się będzie do tychże wielkości wyjściowych. Zwyczajowo definiuje się trzy jakościowe wartości na skali istotności zjawisk – wysoka istotność, średnia oraz niska. Dokonanie rankingu zjawisk w zależności od ich istotności może być wykonane na kilka sposobów, najbardziej rozpowszechnionym z nich jest osąd ekspercki [41], [56], [57]. W ostatnich latach obserwowano również rozwój metod PIRT bazujących na analizach wrażliwości [58], [59], jednak wciąż są to metody innowacyjnie, które nie są szeroko stosowane oraz wdrażane. Korzystając z dostępnych już wyników prac ekspertów którzy dokonywali analizy i rankingu istotności zjawisk zachodzących podczas awarii typu LBLOCA [14], [48], [60] w tabeli 4.1. przedstawiono spis najważniejszych zjawisk w podziale na fazy rozwoju awarii z uwzględnieniem stopnia występowania zjawiska w danej fazie.

Tabela 4.1. Identyfikacja najważniejszych zjawisk zachodzących podczas awarii typu LBLOCA w podziale na fazy rozwoju awarii

Stopień występowania zjawiska: „+” występuje; „o” występuje częściowo; „-” nie występuje			
Zjawisko	Faza rozwoju awarii		
	Wydmuch	Napelnianie	Zalewanie rdzenia
Wypływ przez rozerwanie (wypływ krytyczny)	+	+	+
Wymiana ciepła w rdzeniu	+	+	+
Zmiana charakteru przepływu dwufazowego	o	+	+
Kondensacja oraz mieszanie się pary wodnej i wody	o	+	+
Ograniczanie przepływu przeciwpądowego (CCFL)	o	+	+
Praca pompy w warunkach przepływu dwufazowego oraz jednofazowego w zależności od medium	+	o	o
Bypass wtrysku z systemów zalewania rdzenia reaktora	o	+	o
Propagacja frontu zalewania	o	o	+
Wiązanie pary wodnej	-	o	+
Zachowania gazów niekondensujących	-	o	o

Spśród przedstawionych w tabeli zjawisk jedynie dwa z nich występują w pełni podczas wszystkich faz awarii. Są to wypływ przez rozerwanie oraz zjawiska związane z wymianą

ciepła w rdzeniu reaktora. Wyływ przez rozerwanie określa szybkość ucieczki chłodziwa z obiegu pierwotnego, a więc również spadek ciśnienia w obiegu, co ma bezpośredni wpływ na momenty czasowe w których dojdzie do wyzwolenia sygnałów wtrysku wody z systemów awaryjnego zalewania rdzenia. W związku z występowaniem znaczącej różnicy między ciśnieniem w obiegu a ciśnieniem w obudowie bezpieczeństwa podczas dużego rozerwania rurociągu obiegu pierwotnego obserwuje się wyływ krytyczny. Masowe natężenie wyływu jest wielkością łatwo mierzalną, stąd zjawisko to jest jednym z lepiej zbadanych w instalacjach eksperymentalnych. Drugim zjawiskiem obserwowanym we wszystkich fazach jest wymiana ciepła w rdzeniu. Pod tym dość szerokim pojęciem uwzględniane są zjawiska takie jak występowanie kryzysów wrzenia, krytycznego strumienia cieplnego, zmiany mechanizmów wymiany ciepła czy zapewnienie chłodzenia w warunkach naturalnej oraz wymuszonej konwekcji. Całościowo zjawiska te zazwyczaj charakteryzuje się poprzez badanie temperatury koszulki paliwowej oraz określanie wartości parametru DNBR (ang. „departure from nucleate boiling”) który określa odejście od wrzenia pęcherzykowego, czyli tzw. pierwszy kryzys wrzenia.

Kolejne trzy zjawiska występują częściowo podczas fazy wydmuchu i potem w pełni podczas faz napełniania i zalewania. Wszystkie trzy są powiązane z występowaniem przepływu dwufazowego oraz mieszania się pary wodnej i wody. Są one obserwowane szczególnie w rdzeniu reaktora, szczelinie opadowej, górnej komorze mieszania oraz w rurociągach obiegu pierwotnego i U-rurkach wytwornicy pary. Zjawiska te mają bezpośredni wpływ na efektywność schładzania rdzenia w trakcie rozwoju awarii.

Pięć kolejnych zjawisk występuje jedynie częściowo w niektórych fazach, całościowo jedynie w jednej z faz bądź nie występuje całościowo w żadnej z faz. Zjawiska te również mają znaczący wpływ na przebieg awarii, jednak ze względu na ich ograniczenie w stopniu występowania nie są aż tak intensywnie badane jak zjawiska opisane w dwóch poprzednich grupach.

4.2. Wdrożenie drugiego elementu metodyki

Element drugi metodyki, czyli przygotowanie odpowiedniej dla wybranego scenariusza awaryjnego bazy eksperymentalnej rozpoczyna się od określenia eksperymentów dostępnych dla wybranego problemu. W tabeli 4.1 określono dziesięć najważniejszych zjawisk dla awarii LBLOCA dla których rozpoczęto określanie dostępnych eksperymentów. Eksperymenty obejmują zarówno te w mniejszej skali, w których badane są pojedyncze zjawiska (SETF) oraz eksperymenty przeprowadzone w większej skali w których badane jest wiele zjawisk

jednocześnie (ITF). Na podstawie [14], [15] oraz informacji z bazy danych eksperymentów cieplno-przepływowych dla reaktorów jądrowych (<https://www.oecd-nea.org/tiethysweb/>) w tabeli 4.2 przedstawiono liczbę eksperymentów SETF dostępnych dla danego zjawiska oraz które zjawiska były badane w wybranych instalacjach ITF. Spośród instalacji ITF wybrano cztery które są dostępne w bazie danych OECD NEA oraz które były wykorzystywane do badania zjawisk występujących we wszystkich fazach awarii LBLOCA – LOFT, PKL, SEMISCALE oraz UPTF.

Tabela 4.2. Matryca określająca dostępność instalacji eksperymentalnych dla najważniejszych zjawisk zidentyfikowanych dla awarii LBLOCA

Stopień występowania zjawiska w instalacji ITF: „+” występuje; „o” występuje częściowo; „-” nie występuje					
Zjawisko	Liczba eksperymentów SETF / dostępne w bazie OECD	Instalacje eksperymentalne ITF			
		LOFT	PKL	SEMISCALE	UPTF
Wypływ przez rozerwanie (wypływ krytyczny)	27 / 2	+	+	+	+
Wymiana ciepła w rdzeniu	56 / 6	+	+	o	-
Zmiana charakteru przepływu dwufazowego	52 / 4	+	+	+	+
Kondensacja oraz mieszanie się pary wodnej i wody	12 / 1	+	-	-	+
Ograniczanie przepływu przeciwnieprądowego (CCFL)	32 / 2	o	o	-	+
Praca pompy w warunkach przepływu dwufazowego oraz jednofazowego w zależności od medium	8 / 0	o	o	+	-
Bypass wtrysku z systemów zalewania rdzenia reaktora	17 / 1	+	o	o	+
Propagacja frontu zalewania	36 / 4	+	+	+	-
Wiązanie pary wodnej	0	o	o	o	o
Zachowania gazów niekondensujących	10 / 0	+	-	-	+

Dla zjawiska wiązania pary wodnej w bazie prowadzonej przez OECD nie zidentyfikowano dedykowanych instalacji eksperymentalnych w których badane byłoby to zjawisko. Zjawisko to jest szczególnie istotne w górnej komorze zalewania, ponieważ może powodować ograniczanie wzrostu poziomu wody w rdzeniu w fazie zalewania. W związku z tym jest ono zależne od przebiegu awarii oraz innych zjawisk, które zachodzą wcześniej w jej trakcie. Badane jest więc częściowo w instalacjach w większej skali wraz z innymi zjawiskami zachodzącymi w rdzeniu reaktora. Najwięcej eksperymentów w małej skali dotyczyło trzech

zjawisk związanych z warunkami cieplno-przepływowymi w rdzeniu podczas awarii: wymiana ciepła w rdzeniu, zmiana charakteru przepływu dwufazowego oraz propagacja frontu zalewania. Wymiana ciepła w rdzeniu jest zjawiskiem o najszerszej definicji, stąd eksperymenty SETF dla dwóch pozostałych zjawisk stanowią podzbiory zbioru eksperymentów dla wymiany ciepła. Zjawiska te były również badane w instalacjach LOFT, PKL i SEMISCALE w całości lub częściowo, natomiast w instalacji UPTF badane było jedynie zjawisko zmiany charakteru przepływu dwufazowego i dotyczyło ono w tym wypadku szczeliny opadowej zbiornika reaktora. Dla dwóch kolejnych zjawisk, czyli wypływu przez rozerwanie oraz ograniczania przepływu przeciwpądowego (ang. „countercurrent flow limiting” - CCFL) przeprowadzono około 30 badań, z czego po dwa eksperymenty są dostępne w bazie OECD NEA. Dodatkowo zjawisko wypływu krytycznego zostało zbadane w każdej instalacji w dużej skali, natomiast zjawisko CCFL było częściowo badane w instalacjach LOFT oraz PKL oraz w całości w instalacji UPTF. Dla pozostałych czterech zjawisk liczba wykonanych eksperymentów była już znacznie mniejsza (przedział od 8 do 17), dodatkowo brakuje dostępnych danych z instalacji SETF w bazie OECD NEA opisujących zachowanie gazów niekondensujących oraz prace pomp cyrkulacyjnych. Dla dwóch pozostałych zjawisk dostępne są jedynie pojedyncze dane eksperymentalne. Każde z tych czterech zjawisk jest natomiast odwzorowane w pełnym stopniu w co najmniej jednej instalacji testowej w dużej skali.

Spośród instalacji ITF które pozwalają na ocenę wielu zjawisk podczas jednego eksperymentu w instalacji LOFT siedem z dziesięciu najistotniejszych zjawisk było badanych całościowo, natomiast trzy częściowo. W instalacjach PKL i SEMISCALE po cztery zjawiska były badane w pełni, cztery (PKL) lub trzy (SEMISCALE) częściowo a dwa (PKL) lub trzy (SEMISCALE) nie występowały. W instalacji UPTF sześć zjawisk zostało zbadanych w pełni, jedno częściowo a trzy nie występowały. Kampanie eksperymentalne przeprowadzone w instalacji LOFT są więc najbardziej uniwersalne, pozwalając na pokrycie całego spektrum zjawisk występujących podczas awarii typu LBLOCA we wszystkich fazach jej rozwoju. Po dokonaniu rozszerzenia analizy o sześć mniej istotnych zjawiska, których nie ujęto w tabeli 4.2 okazuje się, że wszystkie te zjawiska były analizowane w całości bądź częściowo podczas każdej z faz awarii tylko w instalacji LOFT.

Kolejnym krokiem dotyczącym przygotowania bazy eksperymentalnej jest weryfikacja jej kompletności oraz reprezentatywności. Kompletność bazy jest w praktyce ściśle związana z dostępnością danych eksperymentalnych. Jak opisano powyżej dla pięciu zjawisk dostępność

eksperymentów SETF w bazie OECD NEA jest mocno ograniczona, stąd lista zjawisk ogranicza się de facto do pięciu pozostałych zjawisk. Próba rozwiązania problemu kompletności mogłoby być poszukiwanie innych źródeł danych eksperymentalnych. Jednak mając na uwadze cel rozprawy doktorskiej uznano, że nie zachodzi taka konieczność, ponieważ wdrożenie metodyki prowadzone jest jedynie w wystarczająco reprezentatywnej części, a nie w całości dla wszystkich najistotniejszych zjawisk. Dla ograniczonego zakresu wdrożenia metodyki kompletność bazy danych jest więc wystarczająca. Dla pięciu pozostałych zjawisk istnieje baza eksperymentalna, o różnej jakości danych, jednak całościowo można uznać, że dane te są wystarczająco adekwatne i reprezentatywne, tak aby była możliwość określenia ostatecznej bazy danych która zostanie wykorzystana w dalszych krokach wdrożenia metodyki.

Spośród pięciu zjawisk dla których baza danych jest adekwatna wybrano zjawiska wypływu przez rozerwanie, wymiany ciepła w rdzeniu oraz propagacji frontu zalewania jako zjawisk dla których przeprowadzony zostanie proces kwantyfikacji niepewności parametrów wejściowych. Zjawiska wymiany ciepła w rdzeniu i propagacji frontu zalewania badane są w tych samych instalacjach eksperymentalnych, stąd można ograniczyć liczbę koniecznych do opracowania modeli instalacji eksperymentalnych nie tracąc przy tym możliwości analizy wielu zjawisk. Dodatkowo te same modele kodu obliczeniowego oraz powiązane z nimi parametry wejściowe kodu mają wpływ na przebieg obu tych zjawisk. Zmiana charakteru przepływu dwufazowego, mimo że badana była w tych samych instalacjach eksperymentalnych to zależna jest od równań do których użytkownik kodu nie ma dostępu nawet z zastosowaniem dostępnych parametrów modelowych kodu, stąd nie będzie przedmiotem analizy w rozprawie. Dwa pozostałe zjawiska, czyli wypływ przez rozerwanie oraz CCFL posiadają zazwyczaj oddzielne specjalne modele empiryczne w kodach obliczeniowych stąd dostępne są parametry modelowe które mogą podlegać kwantyfikacji. Z tych dwóch zjawisk większy wpływ na przebieg awarii ma wypływ przez rozerwanie, stąd zdecydowano się na wybór tego zjawiska do dalszej analizy.

Do wykonania obliczeń wstecznej kwantyfikacji niepewności dla zjawiska wypływu przez rozerwanie wykorzystano wyniki badań w instalacji eksperymentalnej **Marviken** [61]. W ośrodku tym badano wiele pojedynczych zjawisk istotnych z punktu widzenia analiz bezpieczeństwa różnych scenariuszy awaryjnych. Poza kampanią eksperymentalną w której badano wypływ krytyczny przez rozerwanie, badano również między innymi transport aerozoli w obudowie bezpieczeństwa, przecieki z obudowy bezpieczeństwa czy zmiany ciśnienia w obudowie bezpieczeństwa w trakcie fazy wydmuchu podczas awarii LOCA. Kampania eksperymentalna dedykowana wpływowi krytycznemu charakteryzuje się rozległym

zakresem przeprowadzonych testów. Wykonano łącznie 27 testów stosując dysze o różnej długości oraz średnicy. Testy przeprowadzano również ze zróżnicowanymi warunkami początkowymi. Duża liczba oraz zróżnicowanie testów pozwoliło na zastosowanie jednej części testów w kroku jedenastym metodyki a drugiej części w kroku trzynastym. Obliczenia sprawdzające można, więc było wykonywać stosując ten sam model instalacji, jednak dla testów opisanych innymi warunkami początkowymi oraz geometrycznymi. Rozbudowana kampania testowa jak również dostępność dobrej jakości danych były czynnikami którymi kierowano się przy wyborze instalacji Marviken jako źródła danych eksperymentalnych dla zjawiska wypływu krytycznego przez rozerwanie. Bardziej szczegółowy opis instalacji Marviken znajduje się w [rozdziale 5.1](#).

Spośród instalacji w których badano podstawowe zjawiska cieplno-przepływowe w rdzeniu wybrano prostą instalację **FLECHT-SEASET** (Full Length Emergency Core Heat Transfer – Separate Effect and System Effects Test), w której badano wymianę ciepła w fazie zalewania awarii LBLOCA. Przeprowadzono kampanię testową z różnymi warunkami początkowymi oraz stosując rozbudowany system pomiarowy pozwalający na zebranie wysokiej jakości danych doświadczalnych. Paliwo jądrowe było symulowane poprzez zastosowanie grzałek kantalowych. Pęczek grzałek symulujących kasetę paliwową posiadał wysokość zbliżoną do wysokości elementów paliwowych reaktorów typu PWR, stąd uzyskiwano warunki cieplno-przepływowe podobne do tych oczekiwanych w trakcie awarii typu LBLOCA. Podobnie jak dla instalacji Marviken dostępność dużej liczby testów pozwoliła na zastosowanie modelu instalacji do wykonania obliczeń zarówno w kroku jedenastym jak i trzynastym metodyki, co przedstawiono w [rozdziale 5.2](#).

Jako instalacja testowa dla której będzie dokonywana walidacja uzyskanych rozkładów gęstości prawdopodobieństwa w elemencie piątym metodyki wybrana została instalacja **LOFT**. Jest to uzasadnione tym, że jest to instalacja ITF w której w sposób całkowity bądź częściowy badano wszystkie zjawiska zachodzące podczas awarii typu LBLOCA. Zjawiska wypływu przez rozerwanie, wymiany ciepła w rdzeniu oraz propagacji frontu zalewania badane były w pełni. LOFT jest instalacją testową która symuluje działania reaktora typu PWR. W instalacji stosowane jest prawdziwe paliwo uranowe, a maksymalna moc reaktora to 50 MW_t. System pomiarowy dostarcza dokładnych danych na temat warunków cieplno-przepływowych w instalacji. Obiekt składa się z rdzenia reaktora, jednej nitki obiegu pierwotnego z wytwornicą pary, stabilizatorem ciśnienia oraz dwoma pompami cyrkulacyjnymi. Druga nitka obiegu posiada oprzyrządowanie symulujące jedynie wytwornicę pary i pompy, jak również zestaw

zaworów pozwalających na odwzorowanie warunków dużego rozerwania obiegu pierwotnego. W obiekcie zapewniono również awaryjne chłodzenie rdzenia z dwóch akumulatorów oraz systemów wysoko- i niskociśnieniowego wtrysku chłodziwa. W badanym teście L2-5 symulowano całkowite, gilotynowe rozerwanie zimnej nitki obiegu chłodzenia. Dane pomiarowe były zbierane podczas wszystkich trzech faz awarii tj. wydmuchu, napełniania i zalewania rdzenia. Szczegółowy opis instalacji oraz konfiguracji testu L2-5 znajduje się w [rozdziale 6.1](#).

4.3. Wdrożenie trzeciego elementu metodyki

Element trzeci metodyki jest elementem przygotowującym do wykonywania kwantyfikacji niepewności parametrów wejściowych. W tym elemencie określa się kod obliczeniowy który będzie dalej wykorzystywany oraz opracowuje modele instalacji testowych Marviken, Flecht Seaset oraz LOFT. Ostatnim krokiem tego elementu metodyki jest wybór parametrów wejściowych które będą kwantyfikowane w elemencie czwartym oraz pozostałych parametrów które będą wykorzystywane w końcowej analizie niepewności w elementach piątym i szóstym.

Zdecydowano, że do wykonywania wszystkich obliczeń oraz opracowywania modeli matematycznych wybranych instalacji eksperymentalnych posłuży amerykański obliczeniowy kod cieplno-przepływowy TRACE [50], [51] dostępny dzięki porozumieniu między Państwową Agencją Atomistyki a Amerykańską Komisją Dozoru Jądrowego. TRACE został opracowany do wykonywania obliczeń najlepszego szacowania dla awarii typu LOCA oraz innych awarii i stanów przejściowych w reaktorach typu PWR oraz BWR. Kod ten nie służy do wykonywania obliczeń dla ciężkich awarii, ponieważ nie posiada rozwiniętych modeli zachowania paliwa jądrowego po przekroczeniu przez koszulkę paliwową temperatury topnienia. Może być on jednak z powodzeniem stosowany do modelowania zjawisk badanych w instalacjach eksperymentalnych, co stanowi jeden z powodów jego szerokiego stosowania oraz wykorzystania go w niniejszej pracy.

W kodzie TRACE zaimplementowano zestawy cząstkowych równań różniczkowych opisujących przepływ dwufazowy oraz mechanizmy wymiany ciepła. Równania te są rozwiązywane stosując metody numeryczne objętości skończonych. Równania dynamiki płynów rozwiązywane są w jednym wymiarze dla większości komponentów i w trzech wymiarach dla zbiornika reaktora. Sześć podstawowych równań zachowania masy, pędu oraz energii dla cieczy oraz mieszaniny pary oraz gazów uzupełnionych jest o dodatkowe równanie dla gazów niekondensujących oraz zestawu równań domykających. Równania zachowania energii oraz masy są stosowane do środków komórek (nodów) w komponentach, natomiast

równania zachowania pędu dla krawędzi komórek. Stąd wielkości takie jak temperatura, ciśnienie, energia wewnętrzna, gęstość są obliczane dla środków komórek natomiast wielkości takie jak prędkość czy masowe natężenie przepływu obliczane są dla krawędzi komórek.

Podejście do modelowania w kodzie TRACE jest oparte na komponentach, czyli składowych bardziej złożonych systemów. Każdy komponent może zostać podzielony na pomniejsze komórki, czyli nody. Kod posiada specjalne modele komponentów takich jak między innymi rurociągi, stabilizator ciśnienia, pompy, separatory pary, zawory, zbiornik reaktora. Każdy z tych komponentów poza właściwościami standardowymi dla wszystkich komponentów posiada również dedykowane opcje modelowe pozwalające użytkownikowi na możliwość dostosowania modelu danego komponentu pod konkretne zastosowanie. W celu modelowania wymiany ciepła w kodzie wprowadzone są struktury cieplne pozwalające na obliczanie przewodzenia oraz konwekcji w dwóch wymiarach dla kartezjańskiego bądź cylindrycznego układu współrzędnych. Dodatkowo istnieje możliwość stworzenia dedykowanego elementu mającego symulować produkcję energii przez reaktor. Komponenty FILL oraz BREAK pozwalają na odwzorowanie warunków brzegowych w obiegu pierwotnym oraz wtórnym. Ostatnim elementem pozwalającym na wszechstronne modelowanie działania systemów obiektów jądrowych są komponenty kontroli i sterowania, takie jak sygnały, bloki kontrolne czy wyłączniki. Co istotne z punktu widzenia zastosowania metodyki wykonywania obliczeń najlepszego szacowania wraz z oceną niepewności, w kodzie TRACE zaimplementowane są dodatkowe modele wypływu krytycznego oraz wymiany ciepła które charakteryzują się parametryzacją dostępną dla użytkownika.

Krok ósmy polegający na opracowaniu modeli matematycznych każdego z wybranych instalacji eksperymentalnych opisany jest w kolejnych rozdziałach. W przypadku instalacji Marviken jest to [rozdział 5.1](#), w przypadku Flecht Seaset [rozdział 5.2](#) a instalacji LOFT [rozdział 6.1](#). W każdym z tych rozdziałów poza opisem modelu opracowanego w kodzie TRACE znajduje się również podsumowanie procesu kwalifikacji modelu zgodnie z podejściem opisanym w [rozdziale 3.3](#).

Krok dziewiąty, czyli wybór najważniejszych parametrów wejściowych które będą następnie poddane kwantyfikacji zostaje wykonany po przygotowaniu modeli obliczeniowych dla każdej instalacji testowej. Zostanie więc opisany we wskazanych wyżej rozdziałach. Konieczne jest jednak przygotowanie metody obliczeniowej która pozwoli na usystematyzowanie procesu wyboru parametrów wejściowych, tak aby miał on poparcie w ilościowych wskaźnikach. Opracowana została metoda globalnej analizy wrażliwości, która została opisana w [rozdziale](#)

[4.4.](#) Podobnie, aby można było wykonać wszystkie kroki elementu czwartego metodyki konieczne jest opracowanie metody obliczeniowej wykonywania wstecznej kwantyfikacji niepewności. Stąd w [rozdziale 4.5](#) przedstawione zostały podstawy teoretyczne dwóch zaproponowanych metod, które następnie zostaną zaprezentowane w praktyce w [rozdziale piątym](#).

4.4. Metoda globalnej analizy wrażliwości

Proces określania parametrów wejściowych, które zostaną poddane analizie zazwyczaj rozpoczyna się od wykonania lokalnych analiz wrażliwości. W lokalnej analizie wrażliwości przeprowadza się obliczenia dla kilku albo kilkunastu lokalnych wartości z zaproponowanego szerokiego rozkładu wartości parametru. Pozwala to na wstępną analizę wpływu parametru na odpowiedź systemów obiektu jądrowego oraz ocenę czy badana wielkość wyjściowa jest zależna od wybranego parametru wejściowego. Dzięki wykonaniu lokalnej analizy wrażliwości można w prosty i intuicyjny sposób zawężyć wstępną listę parametrów wejściowych. Lokalna analiza nie pozwala jednak na ustalenie istotności parametru, czyli na ocenę wkładu niepewności wynikającej z danego parametru na całkowitą niepewność parametru wyjściowego. Taką analizę wykonuje się stosując metody globalnej wrażliwości [46].

Metody globalnej analizy wrażliwości pozwalają na ocenę istotności parametrów wejściowych poprzez dokonanie rankingu ich istotności. Dodatkowo część metod pozwala na badanie kowariancji między parametrami, co pozwala na identyfikację parametrów, w których jeden z nich jest ściśle zależny od drugiego, co powoduje też, że jego wpływ na niepewność wyjściową, okazuje się być pośredni, a nie bezpośredni. Jednymi z najczęściej stosowanych współczynników pozwalających na ranking istotności parametrów są współczynniki korelacji Pearsona [62] oraz Spearmana [63]. Współczynniki te nie dostarczają informacji dotyczącej relacji przyczynowo skutkowej między dwoma parametrami, ale pozwalają na określenie poziomu powiązania (korelacji) pomiędzy parametrem wejściowym oraz wielkością wyjściową. W przypadku współczynnika Pearsona jest to miara liniowości korelacji natomiast w przypadku współczynnika Spearmana miara monotoniczności relacji. Wnikliwy i szczegółowy opis globalnych analiz wrażliwości znajduje się w [46], gdzie zaprezentowano cztery różne metody wykonywania globalnej analizy wrażliwości. Zaimplementowana w pracy metoda obliczeniowa bazuje na jednej z tych czterech metod – metodzie opartej na analizie wariancji.

Wykorzystywana w opracowanej metodyce metoda globalnej analizy wrażliwości oraz niepewności oparta o analizę Bayesowską została pierwotnie opracowana do wykonywania

analiz dostępności systemów w obiektach przemysłowych bądź jądrowych [64], [65]. W kolejnej pracy [66] po raz pierwszy zastosowano metodę do obliczeń cieplnoprzepływowych w kodzie DARIA (Deterministic Analysis for maRIA reactor) [67] dedykowanym dla reaktora badawczego Maria w Narodowym Centrum Badań Jądrowych w Świerku. W analizie autorzy rozważali trzy parametry wejściowe dla których wykonano analizy niepewności oraz globalną analizę wrażliwości, aby ustalić który z parametrów ma największy wpływ na uzyskiwaną niepewność wielkości wyjściowych. Kolejne wdrożenie metodyki w którym uczestniczył autor rozprawy [68] dotyczyło obliczeń cieplno-przepływowych z wykorzystaniem dwóch kodów komputerowych, kodu systemowego TRACE oraz kodu dedykowanego do analiz ciężkich awarii MELCOR [69]. Dokonano analizy niepewności oraz wrażliwości dla pięciu parametrów wejściowych dla testów wpływu krytycznego w instalacji eksperymentalnej Marviken. Analizy te posłużyły jako punkt wyjścia do implementacji metody w kroku dziewiątym metodyki w rozprawie. Warto nadmienić, że w nazwie metody pojawia się analiza niepewności, co nie jest zawsze ścisłym określeniem i zależne jest od zastosowania metody. W przypadku gdy metoda jest stosowana do wykonywania obliczeń w których nie są dostępne dane eksperymentalne to faktycznie można mówić o analizie niepewności, w innym wypadku jest to w rzeczywistości badanie rozbieżności między danymi eksperymentalnymi a obliczeniowymi, które częściej w literaturze określa się miarą dokładności. Rozróżnienie to pojawiać się będzie w tekście rozprawy, jednakże nazwa metody pozostanie niezmienną.

Metoda łączy w sobie analizę niepewności oraz analizę wrażliwości i wymaga wykorzystania znaczących zasobów obliczeniowych. W celu redukcji liczby obliczeń przy jednoczesnym zachowaniu jakości analizy stosowane są pseudo-losowe sekwencje Sobola małej rozbieżności [70]. Na podstawie wyznaczenia tych sekwencji dokonywane jest losowanie wartości wybranych parametrów z określonych uprzednio zakresów ich zmienności. Całkowita liczba obliczeń (permutacji) zależna jest od czynników takich jak: rząd wrażliwości, liczba parametrów wejściowych **D** oraz oczekiwana liczba próbek **P** pojedynczego parametru. Pierwszy rząd wrażliwości wyznacza zależność wariancji wielkości wyjściowej od pojedynczego parametru wejściowego. Drugi rząd wrażliwości określa wkład do wariancji wielkości wyjściowej wnoszony przez interakcję dwóch parametrów wejściowych. Całkowity rząd wrażliwości określa wpływ wybranego parametru na wariancję wielkości wyjściowej wyznaczony przez pierwszy rząd i wszelkie kombinacje drugiego rzędu. W zależności od rzędu wrażliwości wykorzystując generator wartości losowych Saltelliego[71], oparty o pseudo-losowe sekwencje Sobola uzyskuje się następujące liczby obliczeń:

- Dla obliczeń tylko pierwszego rzędu wrażliwości: $N = P \cdot (D + 2)$;
- Dla obliczeń drugiego rzędu wrażliwości: $N = P \cdot (2 \cdot D + 2)$.

Generator wartości losowych Saltelliego pozwala na wygenerowanie N kombinacji wszystkich wybranych parametrów wejściowych. Dla każdej z tych kombinacji wykonywane są następnie obliczenia z zastosowaniem wybranego kodu obliczeniowego dla modelu instalacji eksperymentalnej. Obliczenia te służą do określenia wariancji uzyskiwanych wyników jako miary ich rozbieżności pomiędzy danymi eksperymentalnymi a wykonanymi obliczeniami. W analizie wyników wygodniejszą miarą rozbieżności jest odchylenie standardowe, które określane jest w tych samych jednostkach co badana wielkość wyjściowa. W analizie wrażliwości, która jest następstwem wykonanych obliczeń dokładności, wykorzystuje się, jednakże wariancję, stąd obie te wielkości są istotne w stosowaniu metody.

W stosowanej metodzie miarą pierwszego rzędu wrażliwości jest indeks Sobola pierwszego rzędu S_1 określany równaniem 4.1[71], [72]:

$$S_1 = S_i = \frac{Var_{X_i}(E_{X_{\sim i}}(Y|X_i))}{Var(Y)} \quad (4.1)$$

Gdzie:

Y – to skalarna wartość wielkości wyjściowej (np. maksymalne natężenie przepływu);

X_i – to i-ty parametr wejściowy modelu;

$X_{\sim i}$ – to macierz wszystkich pozostałych parametrów wejściowych modelu poza i-tym;

$E_{X_{\sim i}}(Y|X_i)$ – to wartość oczekiwana wielkości Y uśredniona wobec możliwych wartości parametrów macierzy $X_{\sim i}$, przy stałej wartości parametru X_i ;

$Var_{X_i}(E_{X_{\sim i}}(Y|X_i))$ – to wariancja obliczana dla wszystkich permutowanych wartości parametru X_i ;

$Var(Y)$ – to wariancja wartości wielkości wyjściowej.

Indeks Sobola pierwszego rzędu jest znormalizowany, więc przyjmuje wartości między 0 a 1, stąd suma wszystkich indeksów Sobola dla pełnego spektrum badanych parametrów wejściowych będzie równa 1.

Miarą drugiego rzędu wrażliwości jest indeks Sobola drugiego rzędu S_2 określany równaniem 4.2:

$$S_2 = S_{i,j} = \frac{Var_{X_{i,j}}(E_{X_{\sim i,j}}(Y|X_i, X_j)) - Var_{X_i} - Var_{X_j}}{Var(Y)} \quad (4.2)$$

gdzie oznaczenia są analogiczne do tych wykorzystywanych do opisanego pierwszego indeksu Sobola, jednakże uwzględniona jest dodatkowo interakcja między parametrami wejściowymi i-tym oraz j-tym. Jest to więc miara wpływu tejże interakcji na wariancję wielkości wyjściowej Y.

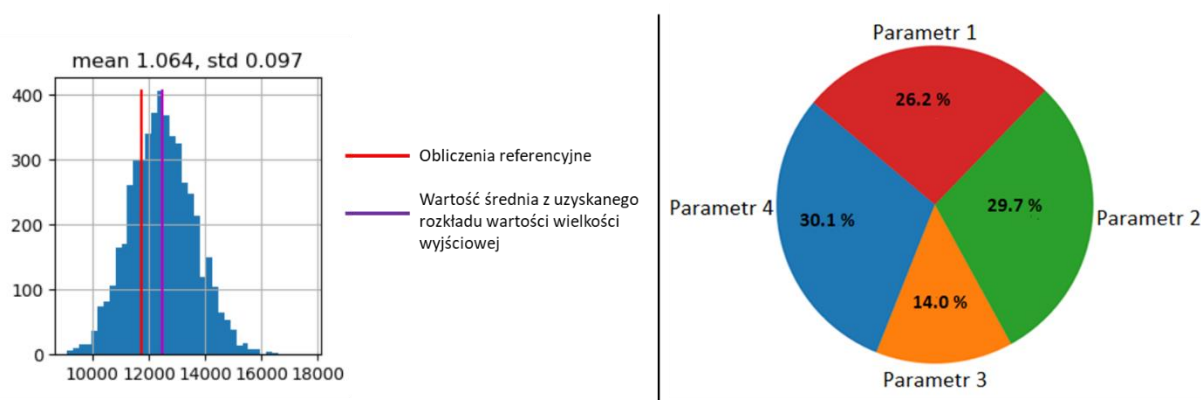
Ostatecznie można zdefiniować całkowity indeks Sobola dla określonego parametru wejściowego zgodnie z równaniem 4.3:

$$S_{Ti} = \frac{E_{X_{\sim i}}(Var_{X_i}(Y|X_{\sim i}))}{Var(Y)} = 1 - \frac{Var_{X_{\sim i}}(E_{X_i}(Y|X_{\sim i}))}{Var(Y)} \quad (4.3)$$

Indeks S_{Ti} określa więc całkowity wpływ danego parametru wejściowego na wariancję wielkości wyjściowej uwzględniając efekty pierwszego rzędu oraz wyższych rzędów. Człon $Var_{X_{\sim i}}(E_{X_i}(Y|X_{\sim i}))$ wyznacza wartość pierwszego rzędu wrażliwości dla macierzy wartości wszystkich parametrów poza i-tym, stąd $Var(Y) - Var_{X_{\sim i}}(E_{X_i}(Y|X_{\sim i}))$ jest miarą wpływu wszystkich składowych wariancji zależnych od i-tego parametru. Właściwości całkowitego indeksu Sobola są istotne między innymi z punktu widzenia minimalizacji złożoności obliczeniowej badanego problemu. Obliczenia miar rzędów wrażliwości wyższych niż pierwszy muszą być wykonywane numerycznie i przy dużym nakładzie obliczeń. Stąd wykonywanie obliczeń dla indeksu Sobola pierwszego rzędu i całkowitego indeksu Sobola pozwala na uzyskanie pełnego spektrum informacji na temat wpływu każdego z parametrów wejściowych, bez konieczności wykonywania obliczeń miar innych rzędów, co znacznie ogranicza czas obliczeń.

Numeryczne algorytmy pozwalające na obliczenia pierwszego oraz całkowitego indeksu Sobola zostały wdrożone stosując metody określone w [71] oraz zaimplementowane w bibliotece SALiB [73] środowiska programistycznego Python [74]. Na rysunku 4.1 zaprezentowano przykład zastosowania opisanej wyżej metody. Na histogramie przedstawiono rozkład uzyskanych rezultatów dla wszystkich obliczeń wrażliwości i zaznaczono wartość referencyjną oraz wartość oczekiwaną z uzyskanego rozkładu. Określona miara rozbieżności, czyli odchylenie standardowe wynosiło 0.097. Na wykresie kołowym zaprezentowano wyniki globalnej analizy wrażliwości w postaci procentowego udziału pierwszego indeksu Sobola dla czterech zmiennych wejściowych.

Metoda może więc być zastosowana zarówno do wyznaczenia najważniejszych parametrów jak i ograniczania ich liczby, tak aby w dalszych analizach wykorzystywać te najbardziej istotne. Wyznaczając całkowity indeks Sobola można też określić łączny wpływ wszystkich rzędów wrażliwości, a więc również interakcji między parametrami, co nie jest możliwe przy zastosowaniu najczęściej wykorzystywanych współczynników korelacji Pearsona i Spearmana.



Rys. 4.1. Przykładowe poglądowe wyniki analizy niepewności (histogram po lewej stronie) oraz globalnej analizy wrażliwości (wykres kołowy po prawej stronie) stosując opisaną metodę

4.5. Metody wstecznej kwantyfikacji niepewności

Opracowanie metody wstecznej kwantyfikacji niepewności jest kluczowym aspektem koniecznym do wdrożenia metodyki oraz było jednym z głównych celów rozprawy. Opierając się na dotychczasowym stanie wiedzy (pokrótce opisanym w [rozdziale 2.1.3](#)) zdecydowano się opracować dwie metody:

1. Opartą stricte na podejściu Bayesowskim i wykonywaniu obliczeń Monte Carlo, jednak bez stosowania metamodeli, których użycie jest obecnie rozpowszechnione w celu minimalizacji liczby obliczeń;
2. Opartą o algorytmy uczenia maszyn, co stanowi nowatorskie podejście, dotychczas nie stosowane w celu wykonywania oceny obliczeń w energetyce jądrowej do kwantyfikacji niepewności wejściowych.

Opracowanie dwóch niezależnych metod pozwoli na ich porównanie i ocenę skuteczności.

4.5.1. Metoda wstecznej kwantyfikacji niepewności oparta na podejściu Bayesowskim

Metoda ta opiera się o wykorzystanie w praktyce twierdzenia Bayesa, które dostarcza informacji na temat prawdopodobieństwa zdarzenia A pod warunkiem, że zajdzie zdarzenie B. W tym celu należy określić prawdopodobieństwo wystąpienia zdarzenia B pod warunkiem zajścia zdarzenia A oraz prawdopodobieństwa wystąpienia niezależnych zdarzeń A oraz B.

W modelowaniu matematycznym oraz naukach fizycznych wnioskowanie Bayesowskie pozwala na aktualizację informacji a'posteriori na temat założeń poczynionych a'priori stosując dane doświadczalne. Założenia a'priori poczynione zostaną wobec rozkładów gęstości prawdopodobieństwa parametrów wejściowych modelu obliczeniowego. Tym samym stosując wnioskowanie Bayesowskie możliwa będzie aktualizacja tych rozkładów uwzględniając rzeczywiste dane z wybranych w kroku szóstym metodyki instalacji eksperymentalnych.

W modelowaniu matematycznym możliwe jest obliczeniowe uzyskiwanie danych $\mathbf{D}$ na podstawie opracowanego modelu $\mathbf{M}$, który opisany jest zestawem parametrów wejściowych Θ [75]. Zapisując $P(D | M, \theta)$, wyznaczone jest prawdopodobieństwo, że uzyskane zostaną dane $\mathbf{D}$ zbieżne z danymi eksperymentalnymi przy zastosowaniu określonego zestawu parametrów wejściowych Θ , modelu matematycznego $\mathbf{M}$. Prawdopodobieństwo to określane jest w statystyce mianem **funkcji wiarygodności** (ang. „likelihood function”). Różne zestawy parametrów wejściowych Θ pozwalają na uzyskanie różnych wartości prawdopodobieństwa, stąd wartości funkcji wiarygodności określają które zestawy parametrów Θ pozwalają osiągnąć największą zbieżność danych obliczeniowych z danymi eksperymentalnymi. Odwracając zależność, otrzymuje się prawdopodobieństwo $P(\theta | D, M)$, określające, że uzyskane parametry modelu $\mathbf{M}$ będą miały wartość z zestawu Θ , przy założeniu, że dane $\mathbf{D}$ są zbieżne z eksperymentalnymi. Jest to więc informacja a'posteriori na temat rozkładu gęstości prawdopodobieństwa zestawu parametrów Θ . Zależność prawdopodobieństwa $P(\theta | D, M)$ od funkcji wiarygodności opisana jest zgodnie z twierdzeniem Bayesa równaniem 4.4:

$$P(\theta | D, M) = \frac{P(D | M, \theta) \cdot P(\theta | M)}{P(D | M)} \quad (4.4)$$

W równaniu tym $P(\theta | M)$ jest funkcją reprezentującą wiedzę (albo jej brak) na temat rozkładu gęstości prawdopodobieństwa parametrów modelu Θ - jest to założenie a'priori. Ostatni człon, czyli $P(D | M)$ reprezentuje prawdopodobieństwo, że model $\mathbf{M}$ dostarcza pełnych informacji na temat danych $\mathbf{D}$ po uwzględnieniu wszystkich możliwych zestawów wartości Θ . Wielkość ta określana jest mianem scałkowanej funkcji wiarygodności (ang. „marginalized likelihood”) ponieważ zgodnie z równaniem 4.5 jest to funkcja wiarygodności scałkowana po wszystkich zestawach wartości Θ .

$$P(D | M) = \int P(D | M, \theta) \cdot P(\theta | M) d\theta \quad (4.5)$$

W modelowaniu matematycznym zakłada się zwyczajowo, że gdy model $\mathbf{M}$ jest poprawny i odpowiednio modeluje dostępne dane $\mathbf{D}$ oraz wnioskowanie Bayesowskie prowadzone jest w celu ustalenia wartości parametrów Θ , to człon ten redukuje się do wartości stałej [75].

Wykonując obliczenia z zastosowaniem modelu matematycznego celem jest jak najlepsze odwzorowanie rzeczywistości. Modele obiektów jądrowych cechują się wysokim stopniem złożoności i zależne są od wielu parametrów, których prawdziwe wartości mogą być znane z określonym błędem pomiarowym. W innym wypadku wartości te są jedynie najlepszym szacowaniem. W związku z tym na potrzeby implementacji wnioskowania Bayesowskiego do metody obliczeniowej wstecznej kwantyfikacji niepewności opisane powyżej dane $\mathbf{D}$ można zapisać jako zestaw wartości $\mathbf{y}_i^M(\mathbf{x}, \Theta)$ gdzie:

- i to i -ta wartość spośród n -kroków czasowych dla wybranej wielkości fizycznej y ;
- M – oznacza, że wartość wielkości fizycznej jest obliczona z zastosowaniem modelu matematycznego (czyli jest rezultatem obliczeń wykonanych z zastosowaniem kodu);
- $\mathbf{x}$ – to wektor wartości wejściowych parametrów projektowych których wartości oczekiwane są mierzalne – są to parametry opisujące warunki początkowe i brzegowe;
- Θ - to wektor wartości parametrów modelowych, takich jak parametry kalibracyjne, współczynniki, stałe modelowe czy stałe fizyczne.

W kontekście wykonywania obliczeń wstecznej kwantyfikacji niepewności warto omówienia są najważniejsze różnice między parametrami wejściowymi należącymi do wektorów $\mathbf{x}$ oraz Θ [27], [76]:

- parametry modelowe Θ występują tylko przy wykonywaniu obliczeń z zastosowaniem modelu, natomiast parametry projektowe $\mathbf{x}$ występują zarówno w modelowaniu jak i przy wykonywaniu fizycznych eksperymentów w instalacjach doświadczalnych;
- parametry modelowe często nie mają faktycznej reprezentacji fizycznej ani nie są mierzalne, stąd ich niepewności nie są znane, w przeciwieństwie do mierzalnych parametrów projektowych których zakresy niepewności można określić poprzez błąd pomiarowy;
- parametry projektowe opisują warunki w teście w danym eksperymencie, są, więc zmienne w zależności od testów i eksperymentów, natomiast parametry modelowe zwyczajowo mają te same wartości bez względu na warunki w eksperymencie;
- rozróżnienie pomiędzy parametrami projektowymi a modelowymi nie jest kluczowe w analizach wrażliwości oraz klasycznej analizie niepewności, gdzie propaguje się niepewności wejściowe przez kod obliczeniowy, jednak jest podstawą wstecznej

kwantyfikacji niepewności, której celem jest określenie rozkładów wartości parametrów modelowych oraz ich wartości oczekiwanych.

Wielkości fizyczne symulowane przez modele matematyczne są więc niezależne od parametrów modelowych Θ i ich faktyczne, rzeczywiste wartości mogą zostać opisane jako $-y_i^R(\mathbf{x})$. W związku z tym, że nie istnieją modele matematyczne które w sposób idealny odwzorowywałyby rzeczywistość, relację między obliczeniami a rzeczywistymi wartościami wielkości fizycznymi określa równanie 4.6:

$$y_i^R(x) = y_i^M(x, \theta) + \delta_i(x) \quad (4.6)$$

gdzie człon $\delta_i(x)$ stanowi rozbieżność pomiędzy wynikami obliczeń (symulacji) oraz rzeczywistymi wartościami wielkości fizycznej. Jest on najczęściej określany jako rozbieżność modelowa (ang. „model discrepancy”) bądź błąd modelowania. Rozbieżność ta jest zawsze obecna w modelowaniu, można ją jedynie minimalizować, ale nie da się jej uniknąć w całości. Reprezentuje ona aproksymacje numeryczne, niepełną wiedzę dotyczącą zależności fizycznych, niedokładność modeli matematycznych oraz niepewność parametrów modelowych.

Dokonując pomiarów wielkości fizycznych uzyskuje się przybliżenie rzeczywistej wartości, obarczonej jednak błędem pomiarowym/błędem obserwacji, co można zapisać równaniem 4.7:

$$y_i^E(x) = y_i^R(x) + \epsilon_i \quad (4.7)$$

gdzie y_i^E to wartość eksperymentalna w i-tym kroku czasowym;

ϵ_i to błąd pomiarowy w tym kroku czasowym.

Ostatecznie można więc powiązać dane eksperymentalne oraz wyniki obliczeń poprzez równanie 4.8:

$$y_i^E(x) = y_i^M(x, \theta) + \delta_i(x) + \epsilon_i \quad (4.8)$$

Celem wstecznej kwantyfikacji niepewności jest uzyskanie rozkładów gęstości prawdopodobieństwa parametrów modelowych Θ . Stosując wnioskowanie Bayesowskie rozwiązaniem problemu są rozkłady a’posteriori określone równaniem 4.4. Jak wspomniano prawdopodobieństwo to jest proporcjonalne do funkcji wiarygodności, która określa w jakim stopniu wyniki obliczeń przy zastosowaniu modelu $\mathbf{M}$ oraz parametrów wejściowych Θ są zbliżone z danymi eksperymentalnymi. Wiążąc funkcję wiarygodności z równaniem 4.8, poszukuje się więc takich wartości Θ , które minimalizują błąd modelowania $\delta_i(x)$. Jest to więc tzw. problem maksymalizacji funkcji wiarygodności. Prawidłowe określenie funkcji

wiarygodności jest, więc kluczowe w poszukiwaniu rozwiązania problemu wstecznej kwantyfikacji niepewności. Wykonując obliczenia z różnymi wartościami parametrów x oraz Θ , stosuje się zawsze ten sam model opisujący instalację eksperymentalną, stąd można założyć stałość modelu M , nie trzeba więc badać wpływu jego zmian na uzyskiwane rezultaty. Stąd rozbieżność pomiędzy danymi eksperymentalnymi a obliczeniowymi wynikać będzie jedynie z dwóch członów opisujących błędy modelowania oraz pomiarowy. Funkcja wiarygodności $P(y_i^M(x) | M, \theta)$ jest więc proporcjonalna do prawdopodobieństw $P(\epsilon_i | M, \theta)$ oraz $P(\delta_i(x) | M, \theta)$. Zakładając [27], [29], [75], [77], że błąd pomiarowy opisany jest rozkładem Gaussa o wartości średniej równej 0 oraz odchyleniu standardowym σ_{mi} , oraz podobnie że błąd modelowy opisany jest rozkładem Gaussa o wartości średniej δ_i oraz odchyleniu standardowym σ_i funkcja wiarygodności może zostać zapisana zgodnie z równaniem 4.9:

$$P(y_i^M) = \frac{1}{\sqrt{2\pi} \sqrt{\sigma_i^2 + \sigma_{mi}^2}} \exp \left\{ \frac{-(y_i^E - y_i^M + \delta_i)^2}{2(\sigma_i^2 + \sigma_{mi}^2)} \right\} \quad (4.9)$$

O ile wartość odchylenia standardowego dla błędu pomiarowego może zostać przyjęta na podstawie danych dotyczących dokładności systemu bądź urządzenia pomiarowego, w przypadku błędu modelowego nie jest to możliwe. Konieczne jest, więc dodatkowe modelowanie tego błędu, które najczęściej przyjmuje formę tzw. metamodelu procesu Gaussowskiego [30], [31], [78] zależnego od czasu. Metamodel to funkcja stosowana w celu aproksymacyjnego opisu relacji pomiędzy danymi wejściowymi a danymi wyjściowymi modelu. Metamodely użyteczne są więc jedynie, gdy istnieje jasna zależność między danymi wejściowymi a wyjściowymi. Podstawowe informacje na temat metamodeli znajdują się w [aneksie](#).

4.5.1.1. Metamodel dla rozbieżności modelowej

Badania przeprowadzane przez autorów w [77], [79] wskazują, że nieuwzględnienie rozbieżności modelowej w analizie bazującej na wnioskowaniu Bayesowskim prowadzi do przesadnego dopasowania (ang. „over-fitting”) wyników obliczeń do danych doświadczalnych co może powodować błędy przewidywania przy wykorzystaniu tych wyników do innych zastosowań. Największą trudnością w prawidłowym modelowaniu rozbieżności modelowej jest jednak brak możliwości jej niezależnego wydzielenia od innych źródeł niepewności, choćby wynikających z niepewności parametrów stanowiących warunki początkowe i brzegowe.

Modelowanie rozbieżności modelowej rozpoczyna się od wydzielenia zbioru kilku eksperymentów o różnych konfiguracjach przeprowadzonych w wybranej instalacji eksperymentalnej. Dla tego wybranego zbioru przeprowadzone zostaną symulacje dla rozbieżności modelowej, natomiast pozostałe eksperymenty zastosowane zostaną do przeprowadzenia wstecznej kwantyfikacji niepewności oraz walidacji. Symulacje rozbieżności modelowej wykonywane są jedynie dla parametrów referencyjnych, czyli z domyślnymi wartościami parametrów wejściowych. Wyniki obliczeń referencyjnych są następnie porównywane z danymi eksperymentalnymi w każdym kroku czasowym eksperymentu. Otrzymuje się więc zestaw wartości $\{y_i^E(x^{ref}) - y_i^M(x^{ref}, \theta^{ref})\}$ dla każdego z i kroków czasowych. W przypadku gdy obliczenia wykonywane są z większą częstotliwością próbkowania to dokonuje się uśrednienia wartości w przedziale czasowym równym krokowi czasowemu z eksperymentu, aby uzyskać jedną wartość $y_i^M(x^{ref}, \theta^{ref})$. Zestawy tych wartości dla każdego kroku czasowego zbierane są dla każdego z n testów wybranych do zbioru, dla którego modelowana będzie rozbieżność modelowa. Otrzymuje się więc zestaw $[\{y_i^E(x^{ref}) - y_i^M(x^{ref}, \theta^{ref})\}_1, \{y_i^E(x^{ref}) - y_i^M(x^{ref}, \theta^{ref})\}_2, \dots, \{y_i^E(x^{ref}) - y_i^M(x^{ref}, \theta^{ref})\}_n]$. Dla każdego kroku czasowego dokonywane jest uśrednianie tych n wartości. Ostatecznie otrzymywany jest, więc zestaw wartości $\{\bar{y}_i^E(x^{ref}) - \bar{y}_i^M(x^{ref}, \theta^{ref})\}$ dla każdego kroku czasowego. Zestaw tych wartości stanowi dane wejściowe dla metamodelu procesu gaussowskiego, który będzie symulował rozbieżność modelową. Zgodnie z opisem [aneksie](#) wybrano funkcję korelacji Materna w celu określenia głównych elementów metamodelu procesu Gaussowskiego. Wdrożenie metamodelu rozbieżności modelowej w praktyce zostanie przedstawione w dalszej części pracy.

4.5.1.2. Metoda MCMC Metropolisa – Hastingsa

Po określeniu funkcji wiarygodności oraz składowych koniecznych do jej obliczenia (uwzględniając metamodel rozbieżności modelowej) pozostaje określenie zakresów i rozkładów gęstości prawdopodobieństwa parametrów Θ dla których wartości funkcji wiarygodności będą wyznaczane. W tym celu wykorzystuje się metody próbkowania, a szczególnie metodę Markov Chain Monte Carlo (MCMC) pozwalającą na generowanie próbek z rozkładu, który w najlepszy sposób będzie zbliżony do poszukiwanego rozkładu a’posteriori [80]. Metody MCMC cieszą się dużą popularnością, ponieważ zapewniają bezpośrednią i intuicyjną metodę próbkowania wartości z nieznanego rozkładu prawdopodobieństwa. Celem jest generacja takich próbek, które będą proporcjonalne do poszukiwanego rozkładu a’posteriori i określenie wartości oczekiwanej z rozkładu. Algorytm

MCMC tworzy łańcuch skorelowanych parametrów $\{\Theta_1, \Theta_2, \dots, \Theta_n\}$, gdzie n to zakładana liczba próbek. Liczba iteracji, którą algorytm spędza wokół jednej wartości jest proporcjonalna do poszukiwanego rozkładu gęstości prawdopodobieństwa. Jest to zdecydowaną zaletą metod MCMC, w porównaniu z klasycznymi metodami próbkowania Monte Carlo które eksplorują przestrzeń próbkowania równomiernie. Wadą algorytmów MCMC jest brak określonego momentu, w którym kolejne łańcuchy oraz próbki powinni przestać być generowane.

Algorytmy MCMC mogą być scharakteryzowane przez różne strategie próbkowania oraz generowania łańcuchów. Jednakże dla każdej z tych strategii trzy składowe są zawsze obecne:

- konieczność określenia pierwszego oszacowania dla wartości początkowych parametrów, dla których będą wyznaczane rozkłady gęstości prawdopodobieństwa,
- rozkład, z którego będą losowane propozycje wartości generowane przez algorytm,
- zasada akceptacji/odrzućania propozycji wartości generowanych przez algorytm.

Jedną z najczęściej wykorzystywanych metod MCMC jest metoda Metropolis-Hastingsa[81], [82], którą wdrożono w niniejszej rozprawie. Spośród trzech składowych każdej z metod MCMC, metoda Metropolis-Hastingsa szczegółowo określa tylko zasadę akceptacji/odrzućania propozycji wartości generowanych przez łańcuch. Dwie pozostałe składowe pozostają w gestii użytkownika.

Pierwsze oszacowanie wartości parametrów Θ_0 może być problematyczne i zdecydowanie zwiększać koszt obliczeniowy, gdy oszacowanie to okaże się znacząco odmienne od wartości z poszukiwanego rozkładu a’posteriori parametrów. W przypadku gdy użytkownik posiada częściową wiedzę na temat rozkładu a’priori, stosuje się wartość oczekiwaną z tego rozkładu. Niemniej jednak, zakłada się, że łańcuch Markova powinien mieć pewną swobodę w odejściu od początkowej wartości pierwszego oszacowania. W związku z tym określoną liczbę początkowych obliczeń odrzuca się, a do analizy wykorzystuje tylko kolejne wartości. Tym samym do końcowych wyników działania algorytmu nie zostaną zaliczone wartości parametrów które pozwoliły na odejście od niepoprawnych pierwszych oszacowań i dojścia do prawidłowego zakresu zmienności parametrów.

Rozkład, z którego będą losowane propozycje wartości parametrów generowane przez łańcuch Markova nazywany jest rozkładem propozycji $Q(\Theta^*|\Theta)$. Wybór typu rozkładu propozycji należy do użytkownika. W pozycji [75] autorzy sugerują stosowanie wielowymiarowego rozkładu normalnego dla zestawu parametrów Θ jako pierwszego oszacowania, które zazwyczaj upraszcza działanie algorytmu MCMC. Przy przyjęciu takiego rozkładu wartość

każdego z parametrów jest losowana z rozkładu normalnego, gdzie wartością średnią jest wartość parametru w poprzednim kroku, natomiast odchylenie standardowe jest na stałe określone procentowo w odniesieniu do wartości średniej.

Ostatnim elementem jest określenie zasady akceptacji/odrzućania propozycji. Jak wspomniano powyżej zastosowana została zasada Metropolis-Hastingsa. Główną ideą algorytmu Metropolis-Hastingsa jest porównanie wartości prawdopodobieństwa z rozkładów a’posteriori dla propozycji i obecnego kroku w łańcuchu Markova (bądź pierwszego oszacowania w pierwszym kroku). Prawdopodobieństwa a’posteriori są nieznane, ale wykorzystując twierdzenie Bayesa można porównywać iloczyny funkcji wiarygodności oraz prawdopodobieństwa a’priori dla propozycji oraz obecnego kroku. Ilozyny te są proporcjonalne do wartości prawdopodobieństwa a’posteriori. Zarówno funkcja wiarygodności jak i prawdopodobieństwo a’priori w przeciwieństwie do rozkładu a’posteriori mogą być bezpośrednio policzone, zgodnie z równaniem 4.9 oraz wiedzą na temat rozkładu a’priori. Dodatkowo w algorytmie Metropolis-Hastingsa należy uwzględnić warunkowe prawdopodobieństwa propozycji oraz bieżącego kroku łańcucha. Gdy warunkowe prawdopodobieństwa są równe tj. $Q(\Theta^*|\Theta) = Q(\Theta|\Theta^*)$, to można te składowe pominąć. Wówczas mówimy o algorytmie Metropolis [81] który obowiązuje dla propozycji o symetrycznym rozkładzie. Symetryczność nie dotyczy w tym przypadku kształtu rozkładu, ale oceny czy ruch z pozycji Θ na pozycję Θ^* jest tak samo prawdopodobny jak ruch z pozycji Θ^* na pozycję Θ . Gdy rozkład propozycji nie jest symetryczny, wówczas składowe muszą zostać uwzględnione a prawdopodobieństwa policzone. Opisane uogólnienie algorytmu Metropolis wprowadził Hastings [82]. Porównanie wartości prawdopodobieństwa z rozkładów a’posteriori określane jest współczynnikiem Metropolis zgodnie z równaniem:

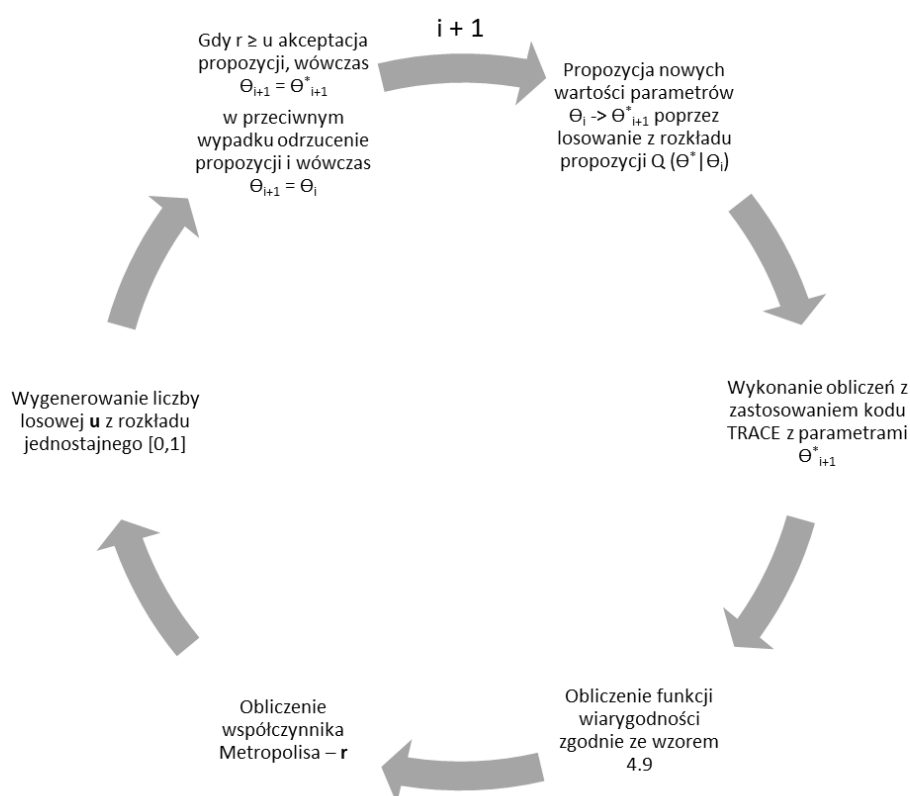
$$r = \frac{P(\theta_{i+1}^*|D, M) \cdot Q(\theta_i|\theta_{i+1}^*)}{P(\theta_i|D, M) \cdot Q(\theta_{i+1}^*|\theta_i)} = \frac{P(D|M, \theta_{i+1}^*) \cdot P(\theta_{i+1}^*|M) \cdot Q(\theta_i|\theta_{i+1}^*)}{P(D|M, \theta_i) \cdot P(\theta_i|M) \cdot Q(\theta_{i+1}^*|\theta_i)} \quad (4.10)$$

Jeśli współczynnik Metropolis jest większy od 1 to propozycja jest zaakceptowana, w przeciwnym wypadku prawdopodobieństwo akceptacji propozycji jest równe wartości współczynnika Metropolis. Prawdopodobieństwo akceptacji propozycji można zapisać:

$$\alpha(\theta_i, \theta_{i+1}^*) = \min(1, r) = \min\left(1, \frac{P(D|M, \theta_{i+1}^*) \cdot P(\theta_{i+1}^*|M) \cdot Q(\theta_i|\theta_{i+1}^*)}{P(D|M, \theta_i) \cdot P(\theta_i|M) \cdot Q(\theta_{i+1}^*|\theta_i)}\right) \quad (4.11)$$

Aby zdecydować, czy propozycja zostanie zaakceptowana, gdy r jest mniejsze od 1, generuje się losową liczbę u z rozkładu jednostajnego ograniczonego w zakresie $[0,1]$. Wówczas, gdy

współczynnik Metropolisa r jest większy lub równy wylosowanej liczbie u , to propozycja jest przyjęta. W przeciwnym wypadku jest ona odrzucona. Algorytm dokonuje więc błędzenia losowego (ang. „random walk”) w przestrzeni rozkładu a’posteriori, akceptując i odrzucając propozycje wartości parametrów w zależności od ich prawdopodobieństwa. Kolejna propozycja jest zależna jedynie od ostatnich wartości parametrów, bez względu na wszelkie pozostałe poprzednie kroki łańcucha. Na rysunku 4.2 przedstawiono schematycznie działanie algorytmu Metropolisa – Hastings zaimplementowane w rozprawie. Algorytm implementujący wszystkie kroki metody napisany został w języku programowania python.



Rys. 4.2. Schemat implementacji algorytmu MCMC Metropolisa-Hastingsa wykorzystywanego do obliczeń wstecznej kwantyfikacji niepewności

4.5.2. Metoda wstecznej kwantyfikacji niepewności oparta o uczenie maszyn

Druga metoda obejmuje nowatorskie zastosowanie algorytmu uczenia maszyn do wstecznej kwantyfikacji niepewności wejściowych parametrów modelowych. Algorytm uczenia maszyn zaimplementowano do oceny miary rozbieżności między danymi eksperymentalnymi a obliczeniami wykonywanymi z zastosowaniem kodu obliczeniowego. Na tej podstawie w kolejnych krokach metody określa się jakie wartości parametrów wejściowych pozwoliły na uzyskanie wyników najbardziej zbliżonych z danymi eksperymentalnymi. Algorytmy uczenia maszyn było wykorzystywane między innymi w ocenie niepewności przy wyznaczaniu

współczynnika mnożenia neutronów [83], [84], przewidywaniu przejścia z wrzenia pęcherzykowego do wrzenia błonowego [85], predykcji wystąpienia wymiany ciepła przy wrzeniu [86] czy przewidywania przebiegu awarii w czasie rzeczywistym [87]. Algorytmy uczenia maszyn wykorzystywano też do wspierania metod wykorzystujących obliczenia MCMC do wstecznej kwantyfikacji niepewności [31], jednak nie odnaleziono pozycji literaturowych w których zaprezentowano by wykorzystanie algorytmów uczenia maszyn bezpośrednio do przeprowadzenia wstecznej kwantyfikacji niepewności.

Jak wykazano w poprzednim rozdziale poszukiwanie rozkładów gęstości prawdopodobieństwa parametrów wejściowych jest silnie zależne od maksymalizacji funkcji wiarygodności, a więc minimalizacji różnicy między danymi eksperymentalnymi, a rezultatami obliczeń. Do obliczeń funkcji wiarygodności określonej równaniem 4.9 wykorzystuje się cały trend czasowy przebiegu zmian. Możliwe jest jednak określenie innych cech pozwalających na porównanie danych eksperymentalnych i rezultatów obliczeń. Kiedy takie cechy przyjmują formę zbioru wartości skalarnych to możliwe jest zastosowanie algorytmów nadzorowanego uczenia maszyn, ponieważ problem porównania jakości danych obliczeniowych w odniesieniu do eksperymentalnych jest podobnie postawiony co problem rozpoznawania wzorców, który jest jednym z podstawowych problemów rozwiązywanych przez uczenie maszyn [88].

Celem zastosowania algorytmu nadzorowanego uczenia maszyn jest poznanie relacji między danymi wejściowymi a wyjściowymi algorytmu. Co warto podkreślić dane wejściowe i wyjściowe algorytmu uczenia maszyn nie są tożsame z danymi wejściowymi i wyjściowymi modelu matematycznego wykorzystywanego do wykonywania obliczeń niepewności dla instalacji eksperymentalnych. Aby określić relację między danymi wejściowymi a wyjściowymi algorytmu konieczne jest przygotowanie zbioru uczącego. Dane wejściowe do algorytmu opisuje się zbiorem cech. Natomiast dane wyjściowe algorytmu reprezentowane są w formie klas, pozwalających na określenie stopnia dopasowania wyników obliczeń do danych eksperymentalnych. Dla zbioru uczącego to użytkownik określa klasy dla obliczeń należących do tego zbioru. Przyporządkowując klasy określa się stopień dopasowania obliczeń do danych eksperymentalnych, co jest odpowiednikiem wykonania jakościowej analizy dokładności obliczeń komputerowych. Następnie po określeniu wartości cech dla wszystkich obliczeń znajdujących się w zbiorze uczącym, algorytm ma za zadanie jak najlepszego „nauczenia się” zależności między wartościami cech (dane wejściowe), a przypisaniem do klas (dane wyjściowe). Posiadając wiedzę na temat relacji między cechami a klasami algorytm otrzymuje w kolejnym kroku zestaw danych tworzących zbiór testowy i ma za zadanie określić klasy dla

tych obliczeń na podstawie ich cech. W tym kroku algorytm nie zna więc prawdziwej relacji między danymi wejściowymi a wyjściowymi zbioru testowego, bazuje na relacji, której nauczył się podczas analizy zbioru uczącego. W opracowanej metodzie wstecznej kwantyfikacji niepewności zestaw cech, który stanowi reprezentację danych wejściowych dla algorytmu uczenia maszyn, określany jest na podstawie wartości istotnych wielkości wyjściowych z obliczeń stanowiących zbiór uczący bądź testowy. Łącząc wartości cech oraz klas, można otrzymać zbiór opisujący zarówno ilościowo jak i jakościowo dokładność obliczeń w odniesieniu do danych eksperymentalnych. Przy określeniu klas podobnie jak w problemach rozpoznawania wzorców [88] możliwe jest więc określenie, które wyniki obliczeń (opisane cechami) są najbardziej zbieżne z danymi eksperymentalnymi. Wartości parametrów wejściowych które pozwolą na uzyskanie wyników zaklasyfikowanych do klasy określającej najlepsze dopasowanie stanowić będą rozwiązanie problemu wstecznej kwantyfikacji niepewności i utworzą rozkłady gęstości prawdopodobieństwa tych parametrów.

4.5.2.1. Cechy

Zbiór cech stanowi reprezentację skłarnych wartości, które w jak najdokładniejszy sposób opisują wielkość wyjściową badaną w eksperymencie. Zgodnie z opisem przedstawionym w rozdziale 4.2 jako badane zjawiska wybrano wpływ przez rozerwanie, wymianę ciepła w rdzeniu oraz propagację frontu zalewania. Najistotniejszym parametrem opisującym wpływ przez rozerwanie będzie oczywiście masowe natężenie przepływu. W przypadku wymiany ciepła będzie to temperatura koszulki, a w przypadku propagacji frontu zalewania – poziom wody w rdzeniu reaktora. Dla każdej z tych trzech wielkości zestaw cech będzie różny, stąd szczegółowo zostaną one opisane w rozdziale 5. Jednakże kilka cech takich jak maksymalna, minimalna, średnia wartość czy odchylenie standardowe w przedziałach kilku sekundowych będzie uniwersalne dla wszystkich wielkości wyjściowych. Cechy powinny w jak najdokładniejszy sposób opisywać przebieg badanej wielkości wejściowej, tak aby uchwycić wszystkie kluczowe aspekty, na podstawie których dokonać można ilościowej analizy wyników obliczeń. Jeśli wielkość wyjściowa reprezentuje zmienne warunki w systemie obiektu jądrowego (np. zmiana charakteru przepływu z przechłodzonego na dwufazowy) to cechy powinny być tak określone, aby pozwalać na uchwycenie tej zmiany. Jest to więc de facto dobór jak najlepszych mierników dla ilościowej analizy wyników obliczeniowych, stąd można posiłkować się danymi literaturowymi [37], [89].

4.5.2.2. Klasy

Klasy stanowią dane wyjściowe z algorytmu uczenia maszyn. W przypadku wykonywania obliczeń niepewności otrzymuje się wyniki symulacji komputerowych i można je porównać do danych eksperymentalnych. Klasy mogą stanowić reprezentację stopnia zgodności między danymi pomiarowymi a wynikami obliczeń. Zgodność tą bada się dla zbioru wielkości wyjściowych, dla których definiowano również cechy, tak aby istniało powiązanie między cechami a klasami. Definicja klas była różna dla każdej z badanych w dwóch eksperymentach wielkości wyjściowych. Istotne jest jednak, aby podobnie jak przy definicji cech, określić klasy w sposób pozwalający na jednoznaczną jakościową ocenę wyników obliczeń, uwzględniając specyfikę każdej z badanych wielkości wyjściowych. Dla wszystkich trzech badanych w rozprawie wielkości wyjściowych zdefiniowano pięć klas. Pozwoliło to zarówno na uproszczenie procesu wykorzystania algorytmu uczenia maszyn dla dwóch różnych eksperymentów, jak również bardziej dokładne porównanie uzyskiwanych rezultatów klasyfikacji. Szczegółowe definicje klas dla każdej analizowanej wielkości wyjściowej znajdują się w rozdziale 5. Ogólnie klasy A – E można jednak scharakteryzować następująco:

- A. Satysfakcjonująca zgodność z danymi eksperymentalnymi w całym czasie trwania eksperymentu;
- B. Satysfakcjonująca zgodność z danymi w jednej części eksperymentu;
- C. Satysfakcjonująca zgodność z danymi w drugiej części eksperymentu,
- D. Brak zgodności z danymi eksperymentalnymi w całym czasie trwania eksperymentu;
- E. Wyniki znacznie odbiegające od przebiegu eksperymentalnego, nierzeczywiste bądź obliczenia przerwane.

Z punktu widzenia wstecznej kwantyfikacji niepewności interesująca jest szczególnie klasa A, jednak pozostałe klasy muszą zostać zdefiniowane, aby w jak najlepszy sposób wyuczyć algorytm, tak aby rozpoznawał różnice między klasami.

4.5.2.3. Zbiór uczący

Zbiór (bądź zestaw, w niniejszej rozprawie sformułowania te są stosowane zamiennie) uczący to w metodzie wstecznej kwantyfikacji niepewności zbiór obliczeń, w którym dla każdego pojedynczego przypadku w zbiorze użytkowników określa klasę, tworząc w ten sposób podstawę do nauczania algorytmu relacji między danymi wejściowymi (cechami) a wyjściowymi (klasami). W uproszczeniu można określić zbiór uczący jako zbiór przykładów przyporządkowujących klasy do zbioru cech. Wybór zbioru uczącego ma kluczowy wpływ na

skuteczność oraz dokładność algorytmu stąd jest to jeden z najważniejszych kroków metody. W wybranych instalacjach eksperymentalnych Marviken oraz Flecht Seaset przeprowadzono serię testów z różnymi warunkami początkowymi oraz brzegowymi, stąd możliwy był podział tych testów tak aby utworzyć zbiory uczący oraz testowy jak również pozostawić testy na cele walidacji. Podziału testów między zbiory dokonywano na podstawie ich grupowania. Testy należące do jednej grupy posiadały zbliżone jakościowo przebiegi analizowanych wielkości wyjściowych. W ten sposób losując te same zestawy parametrów wejściowych dla testów z tej samej grupy oczekiwane były podobne zmiany w przebiegach wielkości wyjściowych wewnątrz grupy. Wielkość zbioru testowego różnić się więc będzie w zależności od liczby grup podobieństwa.

Po określeniu testów które zostaną wykorzystane w formowaniu zbioru uczącego należy określić, ile obliczeń zostanie wykonanych dla każdego z testów. Liczba ta jest zależna od stopnia wiedzy na temat zakresów zmienności parametrów. Im większy stopień niewiedzy na temat zakresów oraz wartości oczekiwanych parametrów tym konieczna jest większa liczba obliczeń, ponieważ zakładane zakresy będą szersze. Przy dużym stopniu niewiedzy należy spodziewać się osiągnięcia przeważającej liczby wyników przyporządkowanych do klasy D, ponieważ kombinacje parametrów będą często zawierały wartości, które będą znacząco odbiegać od wartości oczekiwanych. Przy małej liczbie obliczeń przyporządkowanych do klasy A można wykonać dodatkowe obliczenia z zawężonymi zakresami zmian parametrów. W ten sposób algorytm nie zostanie przetrenowany tylko na jednej klasie – klasie D. Obliczenia wykonuje się stosując metody Monte Carlo, oparte o generatory liczb losowych Saltelliego, opisane w podrozdziale 4.4 bądź klasyczne generatory liczb losowych z zadanego rozkładu dostępne w bibliotece statystycznej języka programowania python.

W algorytmach nadzorowanego uczenia maszyn to użytkownik przypisuje klasę do każdego z wykonanych obliczeń. W opracowanej metodzie odbywa się to w formie analizy jakościowej wyników w odniesieniu do danych eksperymentalnych. Tym samym cechy określone dla każdej wielkości wyjściowej nie są stosowane na tym etapie. Będą one dopiero wykorzystane przy trenowaniu algorytmu po wyborze klasyfikatora. Jakościowa analiza pozwala jednocześnie na określenie dodatkowych cech, ponieważ dokonując podziału między klasami użytkownik musi przyjąć pewne indykatory, które następnie można przełożyć na cechy. Przypisanie klas do każdego pojedynczego obliczenia obarczone może być ludzkim błędem systematycznym wynikającym z subiektywizmu i braku zachowania ciągłości w przypisywaniu podobnych przebiegów do tych samych klas. Aby ograniczyć efekt użytkownika zalecane jest,

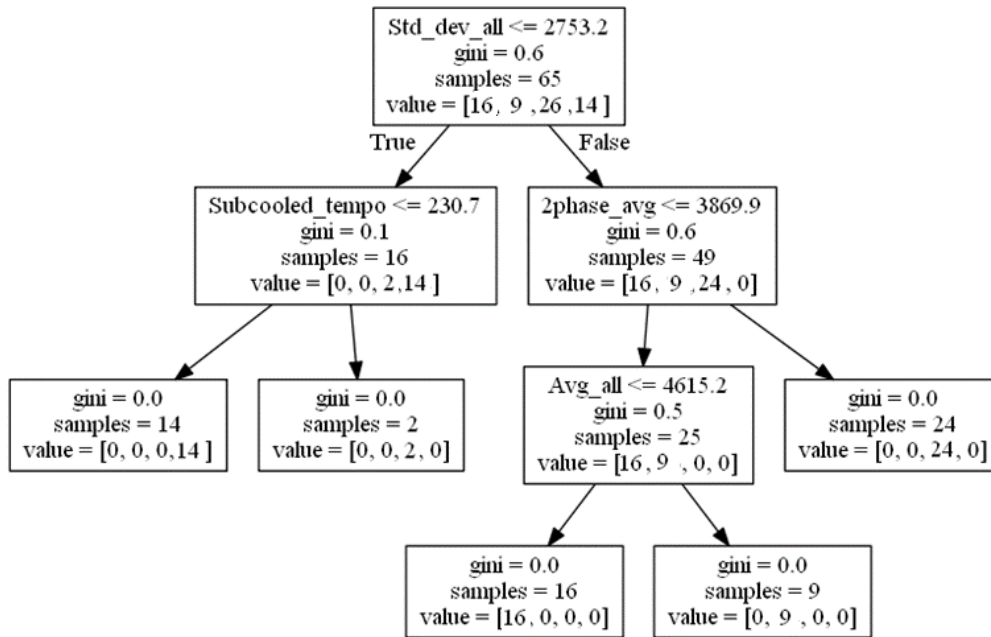
aby stosować proste indykatory opisujące klasy. Stosując algorytmy uczenia maszyn dla celów niniejszej rozprawy zdecydowano, że gdy wartości obliczeń są zbliżone do wartości eksperymentalnych przez więcej niż 60% czasu trwania obliczeń to obliczenia te można zaklasyfikować do klasy A. Podobne podejście stosowano dla klas B oraz C z zastrzeżeniem innych przedziałów czasowych, które są określone w definicji tychże klas. Dla tych klas kluczowe jest również poprawne określenie momentu, w którym kończy się okno czasowe klasy B, a zaczyna się okno czasowe klasy C.

4.5.2.4. Klasyfikator i trenowanie algorytmu

W celu wytrenowania algorytmu, dla którego określono już zbiór uczący, wykonano obliczenia dla zbioru uczącego, przypisano klasy do obliczeń oraz określono wartości cech dla każdego z obliczeń, konieczny jest klasyfikator który będzie dokonywał przypisywania klas do obliczeń w zbiorze testowym. Zdecydowano się wykorzystać klasyfikator lasów losowych, który jest często stosowany w problemach rozpoznawania wzorców [90]. Głównymi zaletami wykorzystania tego klasyfikatora są możliwości wykorzystania go na dużych zbiorach danych, brak konieczności preselekcji wartości wejściowych do algorytmu oraz jego uniwersalność w zastosowaniach do różnych typów danych. Dodatkową zaletą jest możliwość estymacji istotności cech w procesie klasyfikacji.

Klasyfikator lasów losowych opiera się na niezależnych decyzjach podejmowanych przez określoną liczbę drzew decyzyjnych. Ostateczna klasyfikacja jest wynikiem „głosowania” wszystkich drzew decyzyjnych składających się na las losowy. Każde drzewo decyzyjne dokonuje klasyfikacji na podstawie analizy wybranej liczby cech. Liczba ta jest określana przy definiowaniu klasyfikatora lasów losowych. Zgodnie z zaleceniami zawartymi w książce [88] liczba cech stosowana przez pojedyncze drzewo losowe powinna być równa pierwiastkowi kwadratowemu z całkowitej liczby cech. Pojedyncze drzewo decyzyjne nie stosuje wszystkich dostępnych cech, więc stanowi tzw. słaby klasyfikator. Teoria lasów losowych zakłada, że zespół słabych pojedynczych klasyfikatorów tworzy siłę lasu losowego jako klasyfikatora. Na rysunku 4.3 zaprezentowano przykład pojedynczego drzewa decyzyjnego, które podejmuje decyzje na podstawie czterech cech. Drzewo dokonywało oceny 65 obliczeń (samples). Kryterium oceniane było zgodnie z logiką Boole’a jako prawda lub fałsz. Spośród 65 przypadków wartość odchylenia standardowego (cecha std_dev_all) mniejszą od wartości 2753.2 uzyskało 16 przypadków. W 49 przypadkach wartość ta była większa. Pokazuje to, jakie kryteria oceny cech przyjmują drzewa losowe. Wiersz value pokazuje, ile przypadków spełniających dane kryterium zostało zaliczonych do jednej z czterech klas. Spośród 16

przypadków o wartościach odchylenia standardowego mniejszych od 2753.2 dwa przypadki zaliczono dla klasy C, a 14 - do klasy D. Następnie zostały one rozróżnione stosując kolejną cechę – tempo zmiany wartości w trakcie przepływu przechłodzonego (subcooled_tempo). Drzewo dąży do takiego rozdziału względem cechy, aby w końcowych „liściach” drzewa przypadki należały tylko do jednej klasy. Stosując klasyfikator lasów losowych można więc w intuicyjny sposób zrozumieć sposób działania klasyfikatora, i jako że nie działa on na zasadzie „czarnej skrzynki”, dokonywać jego usprawnień.



Rys. 4.3. Przykładowe drzewo decyzyjne stosowane w procesie klasyfikowania przebiegów

W pracy zastosowano funkcje znajdujące się w bibliotece Scikit-learn[91] języka programowania python aby zaimplementować algorytm uczenia maszyn oparty o klasyfikator drzew losowych. Opracowując klasyfikator stosowano głównie domyślne ustawienia dostępnych funkcji. Jedyną zmianą było zwiększenia wagi dla klasy A do pięciu, podczas gdy pozostałe klasy pozostały z domyślnymi wartościami wag równymi jeden. Podczas testów działania algorytmu zaobserwowano, że zwiększenie wartości wagi dla klasy A pozwala na zaklasyfikowanie większej liczby obliczeń do tej klasy. Jak wcześniej wspomniano, to na podstawie obliczeń zaklasyfikowanych do klasy A opracowywane będą rozkłady zmienności parametrów wejściowych, co stanowi rozwiązanie problemu wstecznej kwantyfikacji niepewności.

Po określeniu cech dla wszystkich obliczeń ze zbioru uczącego, przypisaniu im klas oraz określeniu klasyfikatora, algorytm uczenia maszyn oparty o klasyfikator drzew losowych uczył

się relacji między cechami a klasami. W kolejnym kroku wykorzystując tę relację algorytm będzie mógł przypisać klasy do obliczeń ze zbioru testowego, dla których wyznaczone zostaną wartości cech.

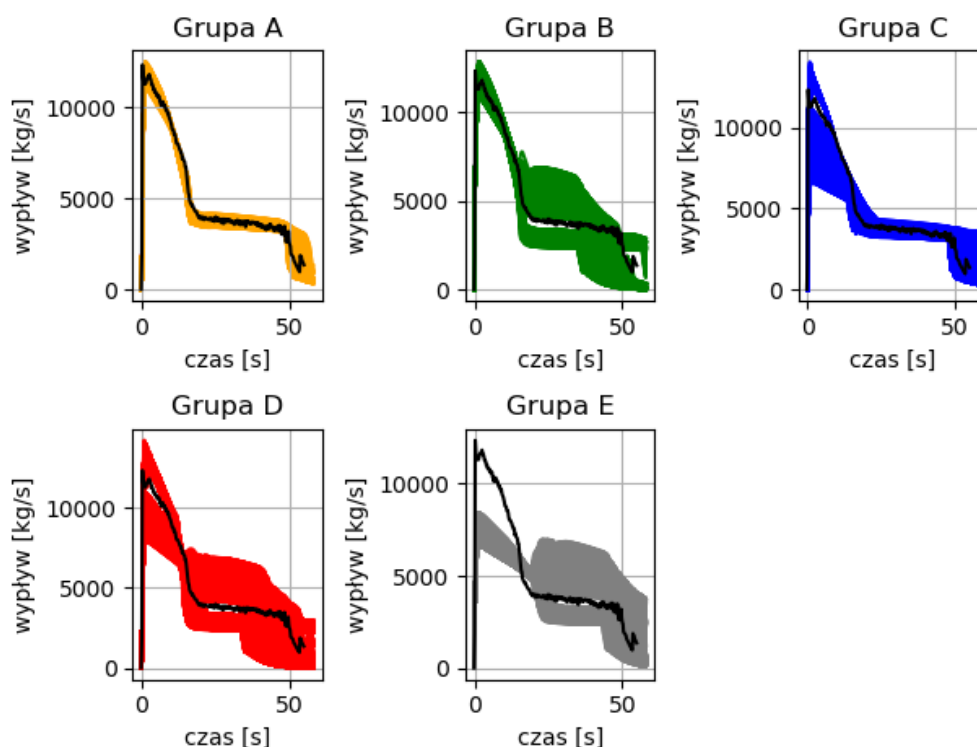
4.5.2.5. Zbiór testowy

Zbiór testowy podobnie jak zbiór uczący powinien stanowić wiarygodną reprezentację warunków panujących w trakcie testów składających się na kampanię testową w danej instalacji eksperymentalnej. Podobnie jak dla zbioru uczącego doboru testów do zbioru testowego dokonywano na podstawie analizy grup podobieństwa między testami. Określenie większej liczby grup i wybór dla każdej z grup co najmniej jednego testu pozwoli na finalne uzyskanie rozkładów gęstości prawdopodobieństwa parametrów wejściowych o większym zakresie stosowalności. Wynika to ze zróżnicowania warunków początkowych oraz brzegowych pomiędzy grupami.

Po określeniu testów które będą tworzyły podstawę przygotowania zbioru testowego, należy określić, ile obliczeń z różnymi wartościami parametrów wejściowych zostanie wykonanych dla każdego testu. Liczba próbek losowanych z rozkładu każdego z parametrów wejściowych powinna wynosić kilka tysięcy. Stosując generator wartości losowych Saltelliego (opisany w rozdziale 4.4) otrzymuje się wówczas kilkadziesiąt tysięcy permutacji wartości parametrów, na podstawie których tworzy się pliki wejściowe do obliczeń. Konieczność wykonania tak dużej liczby obliczeń wynika z dwóch głównych powodów. Po pierwsze wykorzystując algorytm uczenia maszyn konieczne jest uzyskanie wystarczająco dużej liczby obliczeń, które zostaną zaklasyfikowane do klasy A, tak aby wyniki nie były obciążone błędem systematycznym wynikającym ze zbyt małej statystyki. Drugi powód jest związany z wadami dotychczasowych metod analiz niepewności, które opierały się na liczbie obliczeń równej kilkuset przypadkom, co powodowało ich małą powtarzalność. Ostateczna liczba obliczeń jest odmienna w przypadku każdego z analizowanych instalacji testowych, co wynika z różnej złożoności obliczeniowej. W przypadku instalacji Marviken pojedyncze obliczenia trwają kilka sekund, natomiast w przypadku instalacji Flecht Seaset od kilkadziesiąt do kilkuset sekund. Końcowa liczba obliczeń powinna więc być dostosowana do złożoności obliczeniowej oraz dostępnych możliwości obliczeniowych.

Po wyznaczeniu liczby obliczeń należy określić rozkłady gęstości prawdopodobieństwa z których będą losowane wartości parametrów wejściowych. Należy również określić zakresy ich zmienności. Próbkę losowaną są z rozkładów jednorodnych, aby podkreślić brak wiedzy

a’piori na temat rzeczywistych wartości badanych parametrów wejściowych. Zakresy ich zmienności określone są na podstawie wstępnej analizy wrażliwości.



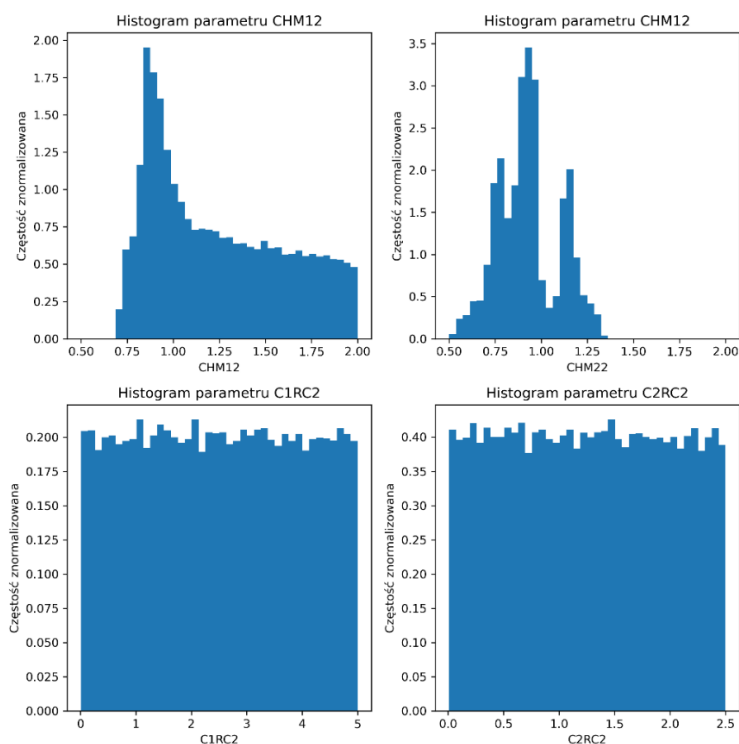
Rys. 4.4. Przykładowe przypisanie obliczeń ze zbioru testowego do pięciu klas przez algorytm uczenia maszyn

W momencie, gdy wszystkie obliczenia dla zbioru testowego zostaną wykonane, można wykorzystać wytrenowany na danych uczących algorytm uczenia maszyn, aby dokonać klasyfikacji obliczeń ze zbioru testowego do jednej ze zdefiniowanych klas stosując opisany klasyfikator lasów losowych. Na rysunku 4.4 zaprezentowano przykładowe przypisanie obliczeń do klas opisanych w rozdziale 4.5.2.2. Można zaobserwować, że algorytm restrykcyjnie przypisywał obliczenia do klasy A, ponieważ rozbieżności pomiędzy przebiegiem eksperymentalnym (czarna linia na rysunku) a przebiegami obliczeniowymi są niewielkie. W grupie B zakres niepewności wyników jest akceptowalny w pierwszej części eksperymentu, a w grupie C w drugiej, co jest zgodne z definicjami tych klas. W grupie D znalazł się największy zbiór obliczeń, które wg. klasyfikacji algorytmu nie osiągnęły wystarczającej zgodności z danymi eksperymentalnymi w którejkolwiek z faz eksperymentu. Przy założeniach rozkładu jednorodnego dla losowanych parametrów, nie jest to zaskoczeniem, ponieważ częściej niż np. w przypadku zastosowaniu rozkładu normalnego będą losowane wartości znacząco odbiegające od rzeczywistego zakresu zmienności parametrów. Uzyskano również liczną reprezentację wyników w klasie E, co pozwala określać zestawy wartości parametrów

które nie powinny być stosowane w przyszłych obliczeniach scenariuszy, w których oczekiwane jest pojawianie się badanego w przykładzie zjawiska.

4.5.2.6. Określenie rozkładów niepewności parametrów wejściowych

Do dalszej analizy, polegającej na określeniu rozkładów niepewności parametrów wejściowych, wykorzystywane są wyniki przypisania do klasy A. Ponieważ w klasie A znajdują się obliczenia o najmniejszej rozbieżności w porównaniu do danych doświadczalnych, to zestawy parametrów które były zmieniane w plikach wsadowych tych obliczeń stanowią zbiór pozwalający na określenie ich rozkładów niepewności. Zestawy wartości parametrów dla wszystkich obliczeń zaklasyfikowanych do klasy A są więc zbierane do oddzielnego pliku. Jest to krok niezależny od samego algorytmu uczenia maszyn, ale wykorzystujący jego działanie. Ze wszystkich wartości dla każdego z parametrów można utworzyć histogramy pokazujące rozkłady ich zmienności oraz zakresy stosowalności. Przykładowe histogramy dla czterech parametrów wejściowych zaprezentowano na rysunku 4.5. Kolejnym krokiem jest zaproponowanie ciągłego rozkładu gęstości prawdopodobieństwa opisującego histogram, tak aby można było z tego rozkładu korzystać w dalszych analizach. Na podstawie zebranych danych można również wyznaczyć wartość najlepszego szacowania. Wartość ta będzie zależna od zaproponowanego rozkładu, w przypadku rozkładu normalnego będzie to wartość średnia, natomiast w przypadku rozkładu log-normalnego będzie to moda ze zbioru wartości.



Rys. 4.5. Przykładowe uzyskane rozkłady niepewności czterech parametrów wejściowych

5. Kwantyfikacja niepewności wejściowych na podstawie obliczeń dla instalacji testowych

W poprzednim rozdziale opisano wdrożenie dwóch pierwszych elementów opracowanej metodyki. W rozdziale 4.3 opisano natomiast element trzeci metodyki, bez demonstracji wykonania prac w kroku ósmym i dziewiątym. Krok ósmy to przygotowanie modeli instalacji eksperymentalnych stosowanych do przeprowadzenia kwantyfikacji niepewności wejściowych dla parametrów które zostaną określone w kroku dziewiątym. Opis modeli dla dwóch wybranych instalacji eksperymentalnych Marviken oraz Flecht-Seaset jest przedstawiony w kolejnych podrozdziałach. Dodatkowo podrozdziały te obejmują opis kroku dziewiątego, czyli wyboru parametrów do analizy stosując metodę BIGUSA dla instalacji Flecht-Seaset oraz lokalną analizę wrażliwości dla instalacji Marviken.

Element czwarty metodyki to przeprowadzenie kwantyfikacji niepewności wejściowych parametrów modelowych stosując modele instalacji eksperymentalnych, opracowane w ramach elementu trzeciego. W tym celu wykorzystywane są dwie metody zaproponowane w niniejszej rozprawie i opisane w [rozdziałach 4.5.1](#) oraz [4.5.2](#). **Jest to główna część obliczeniowa rozprawy, tym samym demonstrująca i porównujące dwie opracowane metody obliczeniowe kwantyfikacji niepewności.** Oczekiwanym rezultatem elementu czwartego metodyki jest określenie rozkładów oraz zakresów zmienności parametrów wejściowych. W ostatnim kroku elementu czwartego wyniki są wstępne sprawdzane poprzez wykonanie obliczeń z wyznaczonymi rozkładami parametrów wejściowych dla wybranych testów wykonanych w instalacjach eksperymentalnych, które nie były wykorzystywane do procesu kwantyfikacji niepewności.

5.1. Obliczenia dla instalacji MARVIKEN

Instalacja eksperymentalna MARVIKEN jest źródłem wielu istotnych danych doświadczalnych wykorzystywanych w badaniu zjawisk fizycznych zachodzących w obiektach jądrowych. W niniejszej pracy wykorzystano dane eksperymentalne dotyczące wpływu krytycznego. W badaniach wpływu krytycznego w instalacji MARVIKEN przeprowadzono 26 testów z różnymi warunkami początkowymi oraz konfiguracjami dysz wylotowych. Tym samym dostępna jest bardzo obszerna baza danych doświadczalnych.

5.1.1. Opis modelowanego zjawiska

Zjawisko wpływu krytycznego jest jednym z istotniejszych zjawisk uwzględnianych w analizach bezpieczeństwa awarii, w których dochodzi do rozerwania rurociągów bądź uszkodzenia zaworów. Jest to jedno ze zjawisk, które jest szczególnie istotne w awarii typu

LBLOCA, ponieważ natężenie wypływu chłodziwa oddziałuje na tempo odkrywania rdzenia oraz spadku ciśnienia w obiegu.

Wypływ krytyczny jest definiowany jako maksymalne masowe natężenie przepływu jakie może być osiągnięte przez ściśliwy płyn przepływający z ośrodka o wysokim ciśnieniu do ośrodka o niskim ciśnieniu. Zjawisko to zachodzi, gdy różnica ciśnień jest większa od ciśnienia krytycznego. Wówczas zmniejszanie wartości ciśnienia w ośrodku o niskim ciśnieniu nie spowoduje zwiększenia przepływu. Wypływ krytyczny będzie więc niezależny od warunków odbiornika. Analiza wypływu krytycznego jest odmienna w warunkach przepływu jedno- i dwufazowego. W warunkach przepływu jednofazowego dla cieczy, wypływ krytyczny jest równy masowemu natężeniu przepływu przy prędkości przepływu równej prędkości dźwięku. W przypadku przepływów dwufazowych prędkość dźwięku zależna jest od objętości frakcji parowej, a sam wydatek przepływu zależy również od temperatury nasycenia.

Z punktu widzenia analizy scenariusza awarii typu LBLOCA istotne jest modelowanie zjawiska wypływu krytycznego przez rozerwanie. Dokładne omówienie zjawiska w tych warunkach można znaleźć między innymi w pracach [92]–[94]. W ostatniej z przywołanych pozycji autorzy dodatkowo opisują podejście do modelowania wypływu krytycznego w systemowych kodach cieplno-przepływowych, do których należy wykorzystywany do obliczeń w tej rozprawie kod TRACE. Strumień przepływu oraz prędkość dźwięku są oceniane w ośrodku o wyższym ciśnieniu, czyli w rurociągu obiegu pierwotnego w przypadku awarii typu LBLOCA. Wykorzystując wbudowane w kod obliczeniowy właściwości płynu, kod dokonuje oceny czy ciśnienie przekroczyło wartość krytyczną w danych warunkach. W większości kodów systemowych stosowane są dodatkowe, specjalne modele empiryczne dla wypływu krytycznego. Wśród popularnych i często wykorzystywanych modeli znajdują się modele Henry-Fasuke[95], Ransom’a i Trapp’a [96] oraz Moody’ego[97].

W kodzie TRACE specjalny model dla obliczeń wypływu krytycznego podzielony jest na trzy pod-modele [51] dla różnych warunków w których może znajdować się płyn. Użytkownik kodu może określić wszystkie miejsca – krawędzie komórek – w których spodziewane jest wystąpienie przepływu krytycznego. Kod obliczeniowy określa w każdym kroku obliczeniowym czy warunki przepływu na wyznaczonych krawędziach komórek spełniają przesłanki, aby uznać przepływ za krytyczny. Jeśli tak, to kod określa w jakim stanie znajduje się płyn i do obliczeń ciśnienia, prędkości oraz masowego natężenia przepływu stosowany jest jeden z trzech pod-modeli wypływu krytycznego:

- Dla cieczy przechłodzonej – stosowany jest zmodyfikowany model Burnell’a do obliczania prędkości cieczy oraz zmodyfikowany model Alamgira [98] do wyznaczenia ciśnienia.
- Dla mieszaniny dwufazowej – stosowany jest model bazujący na modelu Ransom’a i Trapp’a, z dodatkowym uwzględnieniem współczynnika poślizgu.
- Dla nieściśliwego gazu – stosowany jest model oparty o izentropowe rozprężenie gazu idealnego.

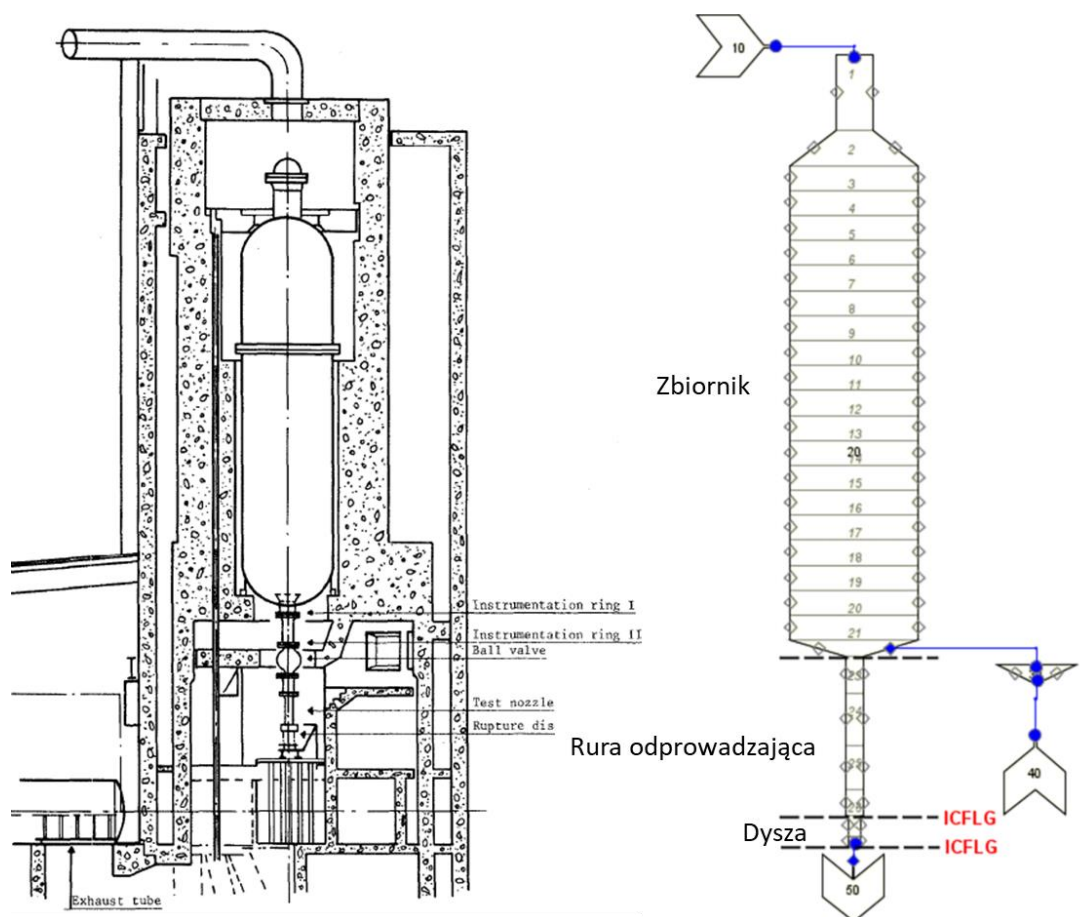
Szczegółowy opis modeli przepływu krytycznego zastosowanych w kodzie TRACE znajduje się w instrukcji kodu [51] oraz w artykule który powstał w trakcie prowadzenia badań nad niniejszą rozprawą [68].

5.1.2. Opis instalacji eksperymentalnej Marviken, kwalifikacja modelu (krok 8)

Konfiguracja instalacji eksperymentalnej Marviken [61] w przeprowadzanych testach przepływu krytycznego obejmowała trzy główne składowe: zbiornik ciśnieniowy o wysokości 24.55 metra i objętości 425 m³, rurę odprowadzającą o długości 6.308 metra i średnicy 0.752 metra oraz dyszę w której znajduje się membrana bezpieczeństwa (ang. „rupture disc”). Dziewięć dysz różniących się pod względem średnicy oraz długości zostało wykorzystanych w dwudziestu siedmiu testach przepływu krytycznego. Testy wykonywane były poprzez wydmuch cieczy bądź mieszaniny wodno-parowej przez dyszę do suchej komory zrzutowej. Momentem rozpoczęcia testu był wzrost ciśnienia do wartości granicznej, przy której rozrywana była membrana bezpieczeństwa. Różnica ciśnienia między zbiornikiem i rurą odprowadzającą a warunkami panującymi w suchej komorze (temperatura pokojowa i ciśnienie równe jednej atmosfery), była na tyle znacząca, że spodziewane było wystąpienie wypływu krytycznego.

Dysze posiadały długość od 166 mm do 1809 mm oraz średnice od 200 do 509 mm. Dodatkowo część dysz nie posiadała jednorodnej powierzchni przepływu na całej długości. W pierwszych czternastu testach stosowano dysze w których średnica membrany bezpieczeństwa była większa od średnicy głównej części dyszy. W testach numer 15, 16, 20, 21, 22 oraz 27 wykorzystywano dyszę w której występowało małe zwężenie w obszarze membrany bezpieczeństwa - z 509 mm do 500 mm. W pozostałych siedmiu testach stosowano dysze o jednorodnym polu powierzchni. W tabeli 5.1 przedstawiono warunki początkowe oraz wymiary dyszy dla każdego z testów, wraz z oznaczeniem, które testy wykorzystywano w obliczeniach dla stosowanych metod kwantyfikacji niepewności.

Na rysunku 5.1 przedstawiono konfigurację instalacji testowej Marviken dla badania wpływu krytycznego oraz opracowany w kodzie obliczeniowym TRACE i środowisku graficznym SNAP model instalacji testowej. W modelu TRACE zbiornik, rura odprowadzająca oraz dysza modelowane są jako jeden komponent PIPE. Zbiornik złożony jest z 22 komórek, rura odprowadzająca z czterech komórek, a dysza z dwóch komórek. Wysokości komórek zostały dobrane w taki sposób, aby ich środki odwzorowywały rozmieszczenie oprzyrządowania pomiarowego, tak aby móc wykorzystać dane pomiarowe dotyczące profilu temperatury i ciśnienia w zbiorniku. U góry zbiornika znajduje się komponent FILL ustalający warunki brzegowe w teście. Wylot z dyszy połączony jest z komponentem BREAK reprezentującym warunki panujące w suchej komorze. Aby zamodelować wystąpienie zjawiska przepływu krytycznego ustawiono flagę ICFLG na granicy między rurą odprowadzającą a dyszą oraz na wyjściu z dyszy.



Rys. 5.1. Instalacja testowa Marviken [61] oraz model instalacji opracowany z zastosowaniem kodu TRACE w środowisku graficznym SNAP

Opracowany model instalacji testowej Marviken został poddany procesowi kwalifikacji. Ze względu na małe skomplikowanie modelu, również sam proces był uproszczony. Wszelkie

wymiary geometryczne (pola powierzchni, średnice hydrauliczne oraz objętości) w modelu są zgodne z wymiarami głównych komponentów instalacji. W stanie ustalonym ciśnienie i temperatura w zbiorniku utrzymywane są na stałym poziomie. Kwalifikacja modelu w warunkach stanu nieustalonego (czyli po rozpoczęciu testu) dla wszystkich wykonanych w instalacji testów przepływu krytycznego została szczegółowo opisana we wspomnianym artykule [68]. Obliczenia referencyjne wykonano dla wszystkich testów z uwzględnieniem zastosowanej dyszy oraz warunków brzegowych i początkowych. Kwalifikacji ilościowej dokonywano poprzez obliczenia błędu względnego między wynikami obliczeń, a danymi eksperymentalnymi dla trzech wybranych punktów w przebiegu masowym natężenia przepływu – maksymalnej wartości natężenia przepływu, jednego punktu w fazie przepływu krytycznego cieczy przechłodzonej (w połowie czasu trwania tej fazy) oraz jednego punktu w trakcie przepływu dwufazowego (w czasie 5 sekund po ustaleniu warunków przepływu dwufazowego). Uśredniony błąd względny dla obliczeń wykonanych dla 26 testów wynosił: 2.8% dla maksymalnej wartości natężenia przepływu, 12.2% dla przepływu krytycznego cieczy przechłodzonej oraz 6.2% dla przepływu dwufazowego. Wyniki te nie odbiegają od wartości przedstawionych w dokumentach walidacyjnych kodu TRACE [99] oraz w opracowaniu porównującym wyniki dla kodów TRACE i RELAP5 dla instalacji Marviken [100]. Warto nadmienić, że przebiegi dla wszystkich testów eksperymentalnych nie mają tego samego kształtu, a maksymalne wartości natężenia przepływu różnią się nawet czterokrotnie. Obserwacja kształtu przebiegów jak również warunków początkowych oraz geometrii dyszy pozwoliła na wyodrębnienie czterech grup podobieństwa dla testów przepływu krytycznego w instalacji Marviken. Przyporządkowanie testów do grup podobieństwa przedstawiono w tabeli 5.1. Grupa pierwsza zawiera osiem testów o średnicy dyszy równej 0.3. Grupa druga zawiera osiem testów, które swym przebiegiem najbliższym przypominają wydmuch chłodziwa w awarii typu LBLOCA i dla których dla obliczeń referencyjnych zaobserwowano najmniejsze wartości błędów względnych. Grupa trzecia zawiera cztery testy, których przebiegi są zbliżone do grupy drugiej, jednak różnią się względem tej grupy stosunkiem długości do średnicy dyszy. Grupa czwarta zawiera sześć testów, których przebiegi znacznie odbiegają od przebiegów charakterystycznych dla pozostałych grup. Przebiegi w grupie czwartej stanowią więc grupę elementów odstających zarówno od pozostałych testów jak i wzajemnie od siebie.

5.1.3. Wybór parametrów wejściowych (kroki 9 oraz 10)

Kod obliczeniowy TRACE pozwala na kalibrację przebiegu przepływu krytycznego poprzez wbudowane cztery parametry modelowe kodu. Pierwsza para parametrów oznaczonych

CHM12 oraz **CHM22** to mnożniki odpowiednio dla przepływu cieczy przechłodzonej oraz mieszaniny dwufazowej. Mnożniki powiązane są z równaniami opisującymi prędkość medium w warunkach wypływu krytycznego. Pozwalają one użytkownikowi kodu obliczeniowego na kalibrację wartości przepływu, aby uwzględnić specyficzne dla danego przypadku zmiany w geometrii dyszy czy rozerwania. Wartości domyślne obu tych współczynników wynoszą 1.0. Druga para parametrów oznaczonych **C1RC2** oraz **C2RC2** to stałe relaksacyjne przepływu krytycznego. Wartości domyślne dla tych parametrów wynoszą odpowiednio 2.0 oraz 1.0.

Ze względu na fakt, że w kodzie dostępne są opisane powyżej parametry które bezpośrednio wpływają na wartość natężenia przepływu nie przeprowadzono dodatkowej analizy wrażliwości pozostałych parametrów modelowych kodu. W artykule [68] przeprowadzono analizę wrażliwości uwzględniając poza parametrami modelowymi, również parametry związane z warunkami początkowymi oraz brzegowymi. Problemem niniejszej rozprawy jest jednak kwantyfikacja niepewności wejściowych parametrów modelowych kodu obliczeniowego, stąd dalsza analiza koncentruje się na określonych powyżej parametrach.

5.1.4. Kwantyfikacja niepewności parametrów wejściowych (krok 11)

Celem kwantyfikacji niepewności parametrów modelowych jest odnalezienie ich wartości najlepszego szacowania oraz ich rozkładu zmienności. Rozkład ten powinien być na tyle uniwersalny, aby przebiegi dla wszystkich testów znajdowały się w wyznaczonym zakresie niepewności.

W celu przeprowadzenia procesu kwantyfikacji niepewności parametrów wejściowych zastosowano dwie opracowane metody opisane w rozdziale czwartym. W obu metodach należy dokonać wyboru testów do wykorzystania na różnych etapach. W przypadku metody opartej o analizę Bayesowską należy dokonać kalibracji metamodelu dla błędu modelowania oraz dokonać obliczeń z zastosowaniem metody Markov Chain Monte Carlo (MCMC). Dla obu tych kroków należy wybrać reprezentatywną liczbę testów. Podobnie w przypadku stosowania metody opartej o uczenie maszyn, gdzie należy wybrać testy będące podstawą do utworzenia zbioru uczącego oraz zbioru testowego. Jak wspomniano powyżej przebiegi można przyporządkować do jednej z czterech grup podobieństwa. Dzięki temu przypisaniu można wybrać reprezentatywne testy dla każdego z kroków obu metod kwantyfikacji niepewności. Przypisanie testów do kroków metod przedstawiono w tabeli 5.1 wraz z warunkami początkowymi i brzegowymi w każdym teście. W tabeli nie uwzględniono testu numer 10, ponieważ ze względu na niesprawność aparatury pomiarowej w trakcie przeprowadzania doświadczenia, nie są dostępne dla niego dane eksperymentalne.

Tabela 5.1. Warunki początkowe oraz brzegowe dla testów przepływu krytycznego w instalacji Marviken wraz z przypisaniem testów do etapów stosowanych metod wstecznej kwantyfikacji

Test	Ciśnienie [MPa]	Temperatura przechłodzenia [°C]	Poziom wody [m]	Długość dyszy L [m]	Średnica dyszy D [m]	L/D [-]	Grupa	Zastosowanie	
								Uczenie maszyn	Bayes
1	5.04	27.3	17.84	1.276	0.3 - 0.4	4.3	1		MCMC
2	4.99	47.7	17.41	1.276	0.3 - 0.4	4.3	1		Błąd modelowy
3	5.03	22.3	17.06	1.976	0.509 - 0.609	3.9	2		
4	4.97	36.2	17.59	1.976	0.509 - 0.609	3.9	2		MCMC
5	4.06	29.1	17.44	1.976	0.509 - 0.609	3.9	2	Zbiór testowy	Błąd modelowy
6	4.96	31.3	17.81	0.671	0.3 - 0.4	2.2	1		MCMC
7	5.01	17.5	17.86	0.671	0.3 - 0.4	2.2	1		
8	4.96	37.5	17.51	1.976	0.509 - 0.609	3.9	2	Zbiór uczący	Błąd modelowy
9	5.05	2.7	18.15	1.976	0.509 - 0.609	3.9	4		
11	4.98	35.5	17.63	1.976	0.509 - 0.609	3.9	2	Zbiór testowy	
12	5.00	32.8	17.52	1.276	0.3 - 0.4	4.3	1	Zbiór testowy	Błąd modelowy
13	5.10	29.5	17.52	0.866	0.2 - 0.4	4.3	1		Błąd modelowy
14	4.98	1.5	18.1	0.866	0.2 - 0.4	4.3	4		
15	5.04	30.4	19.93	1.966	0.509 - 0.5	3.9	2		MCMC
16	5.00	32.9	17.6	1.966	0.509 - 0.5	3.9	2		
17	4.95	29.9	19.85	1.266	0.3	4.2	4		Błąd modelowy
18	5.03	32	17.3	1.266	0.3	4.2	1		
19	5.06	4.4	16.99	1.266	0.3	4.2	4		
20	4.99	6.6	16.65	0.956	0.509 - 0.5	1.9	3	Zbiór testowy	MCMC
21	4.94	33.7	19.95	0.956	0.509 - 0.5	1.9	3		MCMC
22	4.93	50.9	19.64	0.956	0.509 - 0.5	1.9	3		
23	4.96	3.3	19.85	0.391	0.5	0.8	4		Błąd modelowy
24	4.96	32.6	19.88	0.391	0.5	0.8	2		
25	4.92	5.8	19.73	0.661	0.3	2.2	4		Błąd modelowy
26	4.92	33.1	19.31	0.661	0.3	2.2	1		Błąd modelowy
27	4.91	32.3	19.82	0.956	0.509 - 0.5	1.9	3	Zbiór uczący	Błąd modelowy

Dla metody opartej o analizę Bayesowską wybrano dziesięć testów do kalibracji metamodelu dla błędu modelowania. Cztery testy wybrano z grupy pierwszej, trzy testy wybrano z grupy czwartej, dwa testy z grupy drugiej i jeden z grupy trzeciej. Dla kalibracji metamodelu istotne było, żeby w zbiorze znalazły się zarówno testy o dobrym dopasowaniu wyników obliczeń referencyjnych do danych eksperymentalnych jak również o słabym dopasowaniu, aby metamodel faktycznie odzwierciedlał błąd wynikający z modelowania. Do obliczeń MCMC wybrano po dwa testy z grupy pierwszej, drugiej i trzeciej. Nie wybrano testów z grupy czwartej, ponieważ testy z tej grupy nie posiadają cech wspólnych, więc nie było możliwości wyboru reprezentatywnego przypadku. Dodatkowo pozwoli to na ocenę czy zakresy

niepewności uzyskane dla parametrów na podstawie obliczeń dla trzech pozostałych grup pozwolą na pokrycie przebiegów z grupy czwartej przez wyznaczone zakresy niepewności.

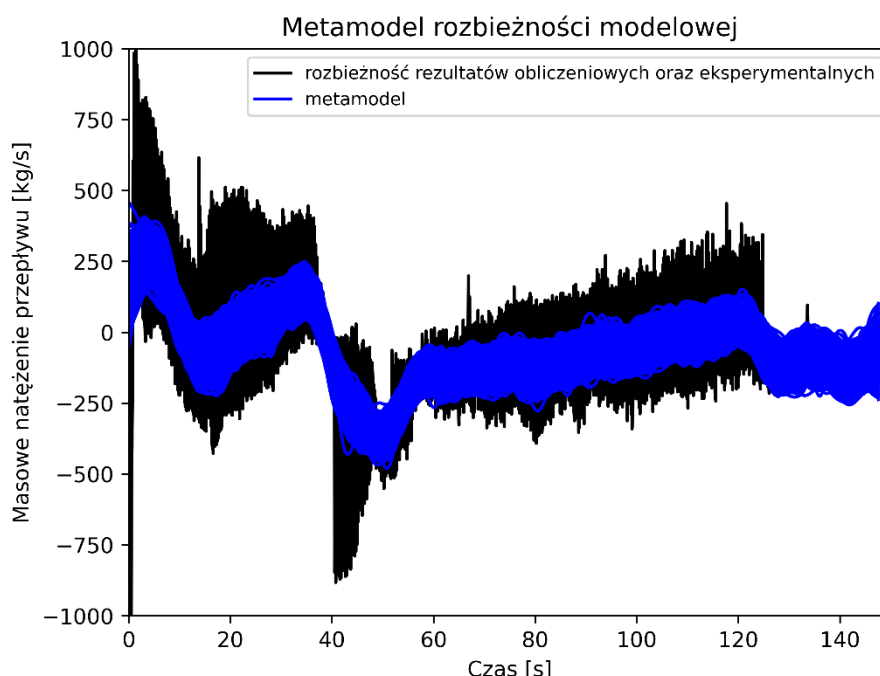
W przypadku metody opartej o uczenie maszyn pierwszym etapem jest wykonanie obliczeń dla zbioru uczącego. Wykonano wstępne obliczenia, aby zaobserwować czy zmiany parametrów, pozwolą na uzyskanie wyników które będzie można przypisać do wielu klas. Pożądane, więc były testy, dla których obliczenia wskazywały na dużą wrażliwość wyników na zmiany parametrów wejściowych, dzięki czemu można łatwiej wyodrębnić większą liczbę testów dla każdej z klas. Początkowo wybrano test numer 8 należący do grupy drugiej jako reprezentatywny dla największej liczby testów. Dzięki wyraźnemu podziałowi na fazę przepływu przechłodzonego oraz dwufazowego, algorytm uczenia maszyn mógł dokładnie rozróżnić te fazy i bardziej dokładnie przypisywać wyniki do klas na podstawie cech. Następnie rozszerzono zbiór uczący o obliczenia dla testu 27 należącego do grupy numer 3. W drugim kroku przypisano cztery testy do stworzenia bazy dla obliczeń zbioru testowego. Zdecydowano się na wybór dwóch testów z grupy drugiej – testu 5 i testu 11, oraz po jednym teście z grupy pierwszej – testu 12 oraz trzeciej – testu 20. Podobnie jak dla metody opartej o analizę Bayesowską nie wykorzystywano testów z grupy czwartej.

5.1.4.1. Kwantyfikacja niepewności parametrów wejściowych stosując metodę opartą o podejście Bayesowskie

Nawiązując do zaprezentowanego w rozdziale 4.5.1 opisu metody kwantyfikacji niepewności opartej o podejście Bayesowskie najistotniejszymi elementami metody, które należy zapewnić przed rozpoczęciem obliczeń wykorzystując metodę MCMC są: metamodel dla rozbieżności modelowej, pierwsze oszacowanie wartości parametrów modelowych, rozkład propozycji, rozkład a’priori parametrów oraz określenie liczby obliczeń i liczby odrzucanych obliczeń.

Metamodel dla rozbieżności modelowej powstał w wyniku porównania wyników obliczeń referencyjnych dla testów oznaczonych w tabeli 5.1 etykietą „Błąd modelowy” z danymi eksperymentalnymi dla każdego z tych testów. Dla każdego z testów zebrano różnice między wynikami obliczeń a danymi eksperymentalnymi w każdym kroku czasowym. Następnie dokonano uśrednienia wszystkich różnic dla każdego kroku czasowego. Ten wektor danych stanowił dane wejściowe do metamodelu, opisanego w [rozdziale 4.5.1.1](#). Dane wyjściowe z metamodelu stanowi zestaw 4000 procesów Gaussowskich reprezentujących przebieg rozbieżności modelowej w czasie. Na rysunku 5.2 zaprezentowano porównanie danych wejściowych oraz wyjściowych z metamodelu. Wartości dla metamodelu mają mniejsze

amplitudy zmian w porównaniu do rezultatów obliczeniowych, jednak ogólny przebieg zmian został zachowany. Metamodel wykorzystywany będzie przy obliczeniach funkcji wiarygodności dla każdej propozycji w obliczeniach z zastosowaniem algorytmu Markov Chain Monte Carlo (MCMC).



Rys. 5.2. Metamodel rozbieżności modelowej dla obliczeń MCMC

Pierwszym oszacowaniem parametrów modelowych CHM12, CHM22, C1RC2 oraz C2RC2 był ich wartości domyślnie, czyli kolejno 1.0, 1.0, 2.0 oraz 1.0. Wyniki obliczeń dla tych parametrów nazywane są obliczeniami referencyjnymi. Wyniki dla pierwszej propozycji wartości parametrów będą porównywane z wynikami dla obliczeń referencyjnych. Jako rozkład propozycji wybrano wielowymiarowy rozkład normalny (liczba wymiarów równa jest liczbie parametrów, w tym wypadku to cztery), gdzie każdy z rozkładów posiadał stałe odchylenie standardowe równe 0.2 oraz wartość średnią przyjmującą wartość parametru w poprzednim kroku. Tym sposobem zapewniona jest symetryczność rozkładu, pozwalająca na pominięcie w obliczeniach współczynnika Metropolis'a (równanie 4.10) członów $Q(\Theta^*|\Theta)$ oraz $Q(\Theta|\Theta^*)$, ponieważ są one sobie równe. Odchylenie standardowe dobrano po wykonaniu wstępnych obliczeń testowych algorytmu, zgodnie z praktyką wartość ta powinna być dobierana przez użytkownika pod konkretne zastosowanie.

Rozkład a priori dla każdego z parametrów przyjęto jako jednorodny w zakresie od 0 do 20. Oznacza to, że wszelkie możliwe wartości w tym przedziale są tak samo prawdopodobne, natomiast wszelkie wartości parametrów wylosowane spoza tego zakresu posiadają

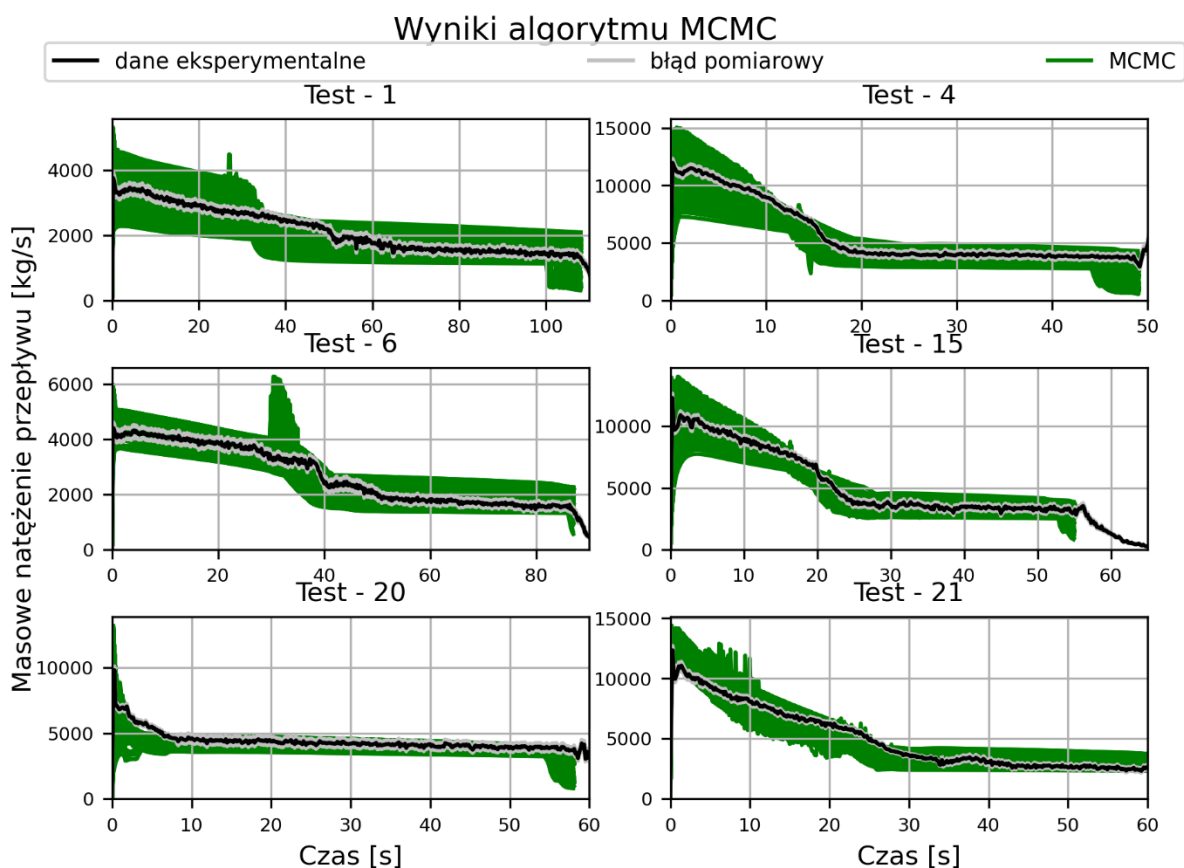
prawdopodobieństwo równe 0. Uwzględniając te prawdopodobieństwa w obliczeniach współczynnika Metropolisa, można zauważyć, że wszelkie propozycje dla parametrów z wartościami powyżej 20 będą odrzucone. Wartość równa 20 została przyjęta arbitralnie jako wartość, która we wstępnych analizach wrażliwości dla wybranych testów wskazywała na uzyskiwanie w obliczeniach bardzo znaczących rozbieżności z danymi eksperymentalnymi. Jednocześnie jest to na tyle duża wartość, że algorytm będzie mógł swobodnie eksplorować obszerną przestrzeń próbkowania.

Dla każdego z testów, zgodnie z tabelą 5.1 przypisanych do wykonania obliczeń MCMC, wykonano 25 000 obliczeń. Łącznie dla sześciu testów wykonano więc 150 000 obliczeń. Obliczenia były wykonywane w ramach implementacji algorytmu opisanego w rozdziale 4.5.1. W pierwszym kroku z rozkładu propozycji losowany jest zestaw nowych wartości parametrów, które są implementowane w pliku wejściowym kodu. Następnie wykonywane są obliczenia z zastosowaniem kodu TRACE. W kolejnym kroku obliczana jest funkcja wiarygodności (równanie 4.9) dla nowej propozycji, z uwzględnieniem niepewności modelowej losowanej ze zbioru 4000 procesów gaussowskich stanowiących metamodel. Wyznaczenie wartości funkcji wiarygodności pozwala na obliczanie współczynnika Metropolisa, którego wartość określa czy dana propozycja wartości parametrów zostanie zaakceptowana. W przypadku gdy wyniki obliczeń dla zaproponowanych parametrów będą bardziej zbliżone do danych eksperymentalnych, to funkcja wiarygodności będzie większa i propozycja zostanie zaakceptowana. Wartości parametrów stanowiących propozycje zostają wówczas zapisane i ostatecznie zbiór tych wartości utworzy rozkład gęstości prawdopodobieństwa a’posteriori dla każdego z parametrów.

Tabela 5.2. Liczba zaakceptowanych propozycji w algorytmie MCMC, wraz z procentowym udziałem zaakceptowanych propozycji w odniesieniu do całkowitej liczby obliczeń

	test 1	test 4	test 6	test 15	test 20	test 21	Suma
Zaakceptowane obliczenia	11 576	4 846	12 846	8 332	4 327	7 485	49 412
Udział procentowy	46%	19%	51%	33%	17%	30%	33%

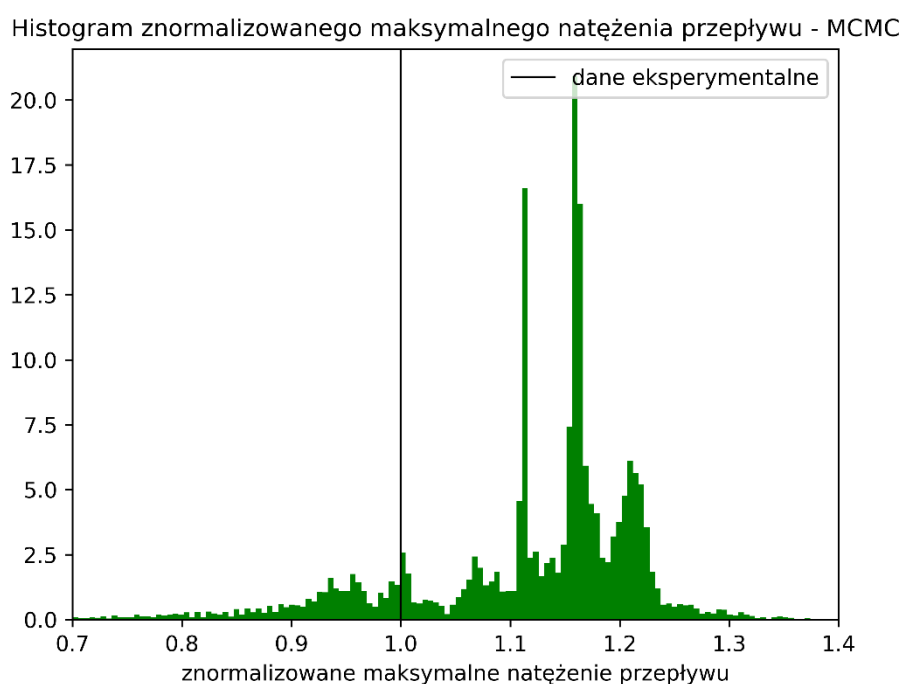
Zgodnie ze standardową praktyką pierwsze 10% propozycji zostało odrzucone z dalszej analizy. W tabeli 5.2 zaprezentowano liczbę obliczeń, które zostały zaakceptowane dla każdego z testów, uwzględniając odrzucenie pierwszych 2500 propozycji. Najwięcej obliczeń algorytm zaakceptował dla testów 6 oraz 1, natomiast najmniej dla testów 20 oraz 4. Łącznie zaakceptowane zostało 33% wszystkich propozycji.



Rys. 5.3. Wyniki obliczeń dla zaakceptowanych przez algorytm MCMC propozycji wartości parametrów dla każdego z testów, w porównaniu z danymi eksperymentalnymi

Na rysunku 5.3 zaprezentowano przebiegi masowego natężenia przepływu dla wszystkich zaakceptowanych propozycji parametrów modelowych. Wyniki obliczeniowe dla porównania zestawiono na tle danych eksperymentalnych. Powstałe ze zbiorów obliczeń zakresy w każdym przypadku obejmują dane eksperymentalne. Wskazuje to na prawidłowe działanie algorytmu, odpowiednio wyznaczającego funkcję wiarygodności i akceptującego propozycje. Należy pamiętać, że w przypadku, gdy współczynnik Metropolis'a r jest mniejszy od 1, dokonywane jest porównanie jego wartości z losowo wybraną liczbą z rozkładu jednorodnego o zakresie od 0 do 1. Stąd może się zdarzyć, że gdy wylosowana z tego zakresu liczba będzie bliska zeru to nawet wyniki o dużej rozbieżności z danymi eksperymentalnymi mogą zostać zaakceptowane. Jest to świadomie wbudowana cecha algorytmu, tak aby nie „spędzał” on zbyt dużo czasu w jednym wycinku rozkładu gęstości prawdopodobieństwa. Jednocześnie powoduje to, że część zaakceptowanych obliczeń odbiega od danych eksperymentalnych, poszerzając uzyskany zakres niepewności przebiegu zmiennej wyjściowej.

W celu przeprowadzenia ilościowej analizy otrzymanych rezultatów działania algorytmu porównano wartości maksymalnego natężenia przepływu dla każdego z zaakceptowanych obliczeń w odniesieniu do wartości maksymalnego wydatku masowego z danych eksperymentalnych dla każdego testu. Zbierano więc znormalizowane wartości $\frac{\max(m_i^{test})}{\max(m_{exp}^{test})}$ ze wszystkich testów zgodnie z tabelą 5.2. Histogram wartości przedstawiono na rysunku 5.4.



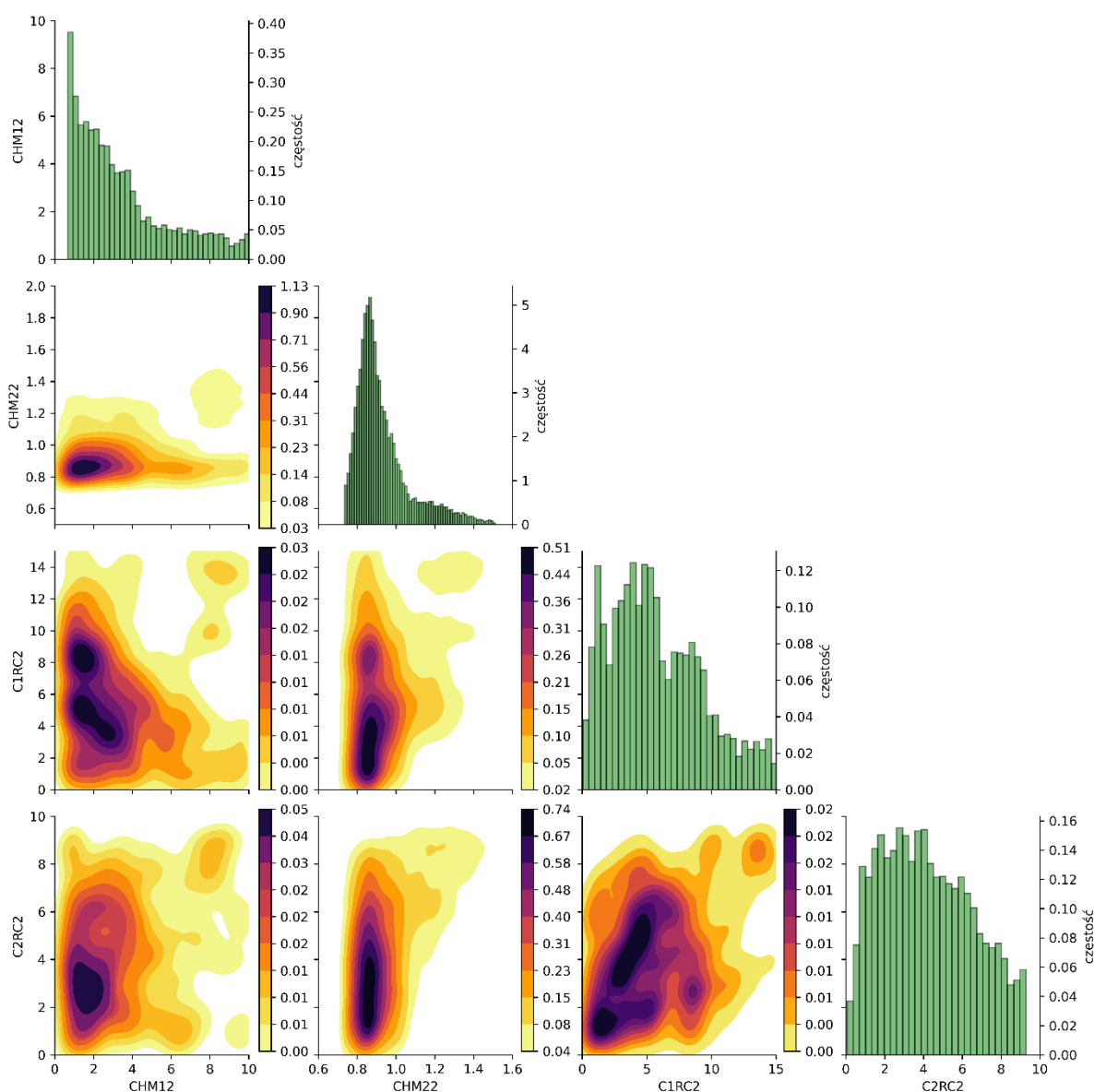
Rys. 5.4. Znormalizowany wydatek przepływu dla wszystkich zaakceptowanych propozycji algorytmu MCMC dla testów 1, 4, 6, 15, 20, 21 w instalacji Marviken

Oszacowane maksymalne natężenie przepływu było w większości przypadków przeszacowane o 10 do 25 procent. Należy pamiętać, że algorytm obliczając funkcję wiarygodności porównuje wartości obliczeń z danymi eksperymentalnymi w każdym kroku czasowym przebiegu, oraz że w obliczeniach uwzględniany był człon rozbieżności modelowych. Uzyskane rozbieżności w wynikach maksymalnego natężenia przepływu, nie są więc znaczące. Jednocześnie jako prosta miara ilościowa będą mogły zostać porównane z rezultatami otrzymanymi dla drugiej metody.

Określenie rozkładów gęstości prawdopodobieństwa parametrów

W celu wyznaczenia rozkładów gęstości prawdopodobieństwa dokonano obróbki parametrów charakteryzujących każdą z propozycji dla sześciu testów, dla których wykonano obliczenia. Dla każdego z czterech parametrów określono minimalne i maksymalne wartości. Następnie zawężono analizę do zakresu od drugiej największej wartości ze zbioru minimów do drugiej

najmniejszej ze zbioru maksimów. W ten sposób wykonano w praktyce operację iloczynu logicznego na rozkładach, aby nie otrzymywać wartości, które dla większości testów będą znacząco wykraczać poza zakresy, w których gwarantują one uzyskanie poprawnych wyników. Jednocześnie, dzięki wykorzystaniu do określenia wartości parametrów sześciu testów o różnych warunkach początkowych i brzegowych, rozkłady te uwzględnią szersze spektrum wartości parametrów, co pozwoli na ich uniwersalne stosowanie.



Rys. 5.5. Jedno i dwuwymiarowe rozkłady gęstości prawdopodobieństwa dla parametrów modelowych CHM12, CHM22, C1RC2, C2RC uzyskane stosując metodę wstecznej kwantyfikacji niepewności opartej o analizę Bayesowską

Ostatecznie na podstawie 11754 zaakceptowanych zestawów propozycji wyznaczono rozkłady gęstości prawdopodobieństwa dla każdego z czterech parametrów. Histogramy rozkładów

gęstości prawdopodobieństwa dla parametrów CHM12, CHM22, C1RC2 oraz C2RC2 przedstawiono na rysunku 5.5 jako ostatni rysunek w każdym z rzędów. Dla histogramów oś y przedstawiona jest po prawej stronie wykresu. Oś y umiejscowiona po lewej stronie stosowana jest dla dwuwymiarowych rozkładów par parametrów. Dwuwymiarowe rozkłady pokazują zakresy wartości w których prawdopodobieństwo uzyskania wyników obliczeń zbliżonych do danych eksperymentalnych jest największe.

Zarówno histogram jak i dwuwymiarowe rozkłady uwzględniające parametr CHM12 wskazują, że zakres wartości między 0.5 a 4 daje największe prawdopodobieństwo uzyskania wyników obliczeń zbliżonych z danymi eksperymentalnymi. Zakres odcięcia prawdopodobnych wartości wynosi 10. Jest to więc szeroki zakres obejmujący wartość domyślną. Zakres z największą gęstością prawdopodobieństwa dla parametru CHM22 to 0.75 – 1.1, natomiast wszystkie wyznaczone wartości mieszczą się w przedziale od 0.7 do 1.5. Rekomendowane jest więc nie stosowanie wartości tego parametru wykraczających poza ten zakres. Jest to jedyny parametr spośród badanych, dla którego można określić bezsprzeczne i klarowne granice rozkładu.

Parametry C1RC2 oraz C2RC2 posiadają podobne funkcje gęstości prawdopodobieństwa a ich najbardziej prawdopodobne wartości znajdują się w zakresie od 0 do 10. Dla parametru C1RC2 wyznaczone zostały wartości ograniczone do 15, jednak są one znacznie mniej prawdopodobne niż wartości w zakresie od 0 do 10. Dla parametru C2RC2 wartości powyżej 10 zostały odcięte jako bardzo mało prawdopodobne.

5.1.4.2. Kwantyfikacja niepewności parametrów wejściowych stosując metodę opartą o uczenie maszyn

Przeprowadzono również kwantyfikację niepewności wejściowych stosując drugą z metod opisanych w rozdziale 4.5 – metodę opartą o algorytm uczenia maszyn. Zgodnie z opisem metody rozpoczęto od określenia cech, klas oraz określenia zbioru uczącego i przeprowadzenia dla niego obliczeń.

Cechy dla obliczeń masowego natężenia przepływu – instalacja testowa Marviken

Jak wspomniano w [rozdziale 4.5.2.1](#) cechy powinny stanowić reprezentatywne mierniki dla analizy ilościowej wyników obliczeniowych w odniesieniu do danych eksperymentalnych. Ze względu na złożoność obliczeniową preferowane jest określanie cech, które są wartościami skalarnymi. Dla przebiegu masowego natężenia przepływu określono łącznie trzynaście cech w podziale na cechy całosciowe (dotyczące całego przebiegu), cechy dla przepływu cieczy przechłodzonej oraz cechy dla przepływu dwufazowego.

Jako całościowe cechy wybrano maksymalne natężenie przepływu, minimalne natężenie przepływu, wartość średnią natężenia przepływu oraz odchylenie standardowe obliczane na podstawie średnich wartości z przedziałów trzysekundowych. Wartość odchylenia standardowego obliczana jest wedle równania 5.1:

$$\sigma = \sqrt{\frac{\sum_{i=1}^n (x_i - Y_i)^2}{n}} \quad (5.1)$$

Gdzie n to liczba trzysekundowych przedziałów w całym czasie obliczeń, i to indeks przedziału, x_i to wartość średnia wydatku przepływu z i -tego przedziału trzysekundowego dla obliczeń natomiast Y_i to wartość średnia wydatku przepływu z i -tego przedziału trzysekundowego dla danych eksperymentalnych.

Dla przepływu cieczy przechłodzonej określono następujące cztery cechy: czas t_1 w którym zakończył się przepływ jednofazowy, wartość średnia natężenia przepływu w czasie t_1 , odchylenie standardowe wyznaczone wedle równania 5.1 jednak tylko dla wypływu przechłodzonego, oraz tempo zmiany natężenia przepływu w czasie t_1 . Czas t_1 wyznaczono na podstawie oszacowania, kiedy temperatura cieczy na dole zbiornika osiąga temperaturę nasycenia. Tempo zmiany wartości natężenia przepływu w czasie t_1 obliczane jest wedle równania 5.2

$$\Delta \dot{m}_{sub} = \frac{\dot{m}_{sub_1} - \dot{m}_{sub_n}}{(n-1) \cdot 3} \quad (5.2)$$

Gdzie n to liczba trzysekundowych przedziałów w czasie t_1 , $\dot{m}_{sub_1}$ oraz $\dot{m}_{sub_n}$ to średnie wartości odpowiednio w pierwszym i ostatnim przedziale trzysekundowych w trakcie trwania przepływu przechłodzonego.

Dla przepływu dwufazowego określono pięć cech: czas t_2 określający moment rozpoczęcia przepływu dwufazowego, czas t_3 określający moment zakończenia przepływu dwufazowego, wartość średnią wydatku przepływu w czasie $t_3 - t_2$, oraz tempo zmian wartości natężenia przepływu obliczane analogicznie jak dla przepływu cieczy przechłodzonej oraz odchylenie standardowe wyznaczone wedle równania 5.1 jednak tylko dla przepływu dwufazowego.

Czasy t_2 oraz t_3 wyznaczano na podstawie opracowanego algorytmu szczegółowo opisanego w artykule [101] którego autor rozprawy jest współautorem.

Klasy dla obliczeń masowego natężenia przepływu – instalacja testowa Marviken

Klasy stanowiąc będą dane wyjściowe z algorytmu uczenia maszyn jako reprezentacja stopnia zgodności między danymi pomiarowymi a wynikami obliczeń. W przypadku wykonywania obliczeń dla instalacji Marviken i badania masowego natężenia przepływu zdefiniowano pięć klas o oznaczeniach od A – E:

- A. Satysfakcjonująca zgodność z danymi eksperymentalnymi w całym czasie trwania eksperymentu;
- B. Satysfakcjonująca zgodność z danymi eksperymentalnymi w fazie wypływu przechłodzonego;
- C. Satysfakcjonująca zgodność z danymi eksperymentalnymi w fazie wypływu dwufazowego;
- D. Brak zgodności z danymi eksperymentalnymi w całym czasie trwania eksperymentu;
- E. Wyniki znacznie odbiegające od przebiegu eksperymentalnego bądź obliczenia przerwane przez wystąpienia błędów.

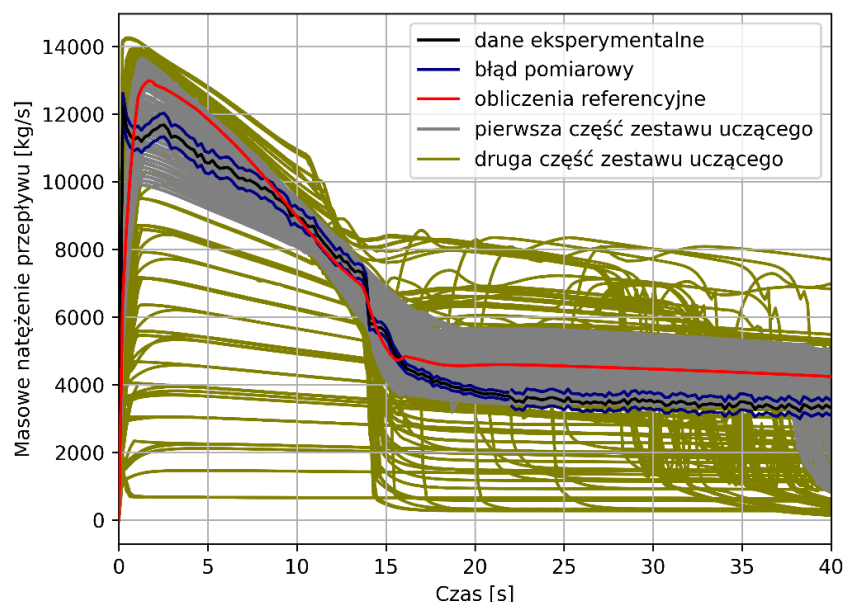
Zbiór uczący dla obliczeń masowego natężenia przepływu – instalacja testowa Marviken

W toku implementacji metody wstecznej kwantyfikacji niepewności stosując algorytm uczenia maszyn przygotowano dwa zestawy uczące. Pierwszy – podstawowy zbiór uczący składał się z obliczeń tylko dla jednego testu. Zgodnie z opisem zamieszczonym dla tabeli 5.1 podstawowy zbiór uczący stanowiły obliczenia dla testu 8. Dla zbioru wykonano łącznie 720 obliczeń. Pierwsze 510 obliczeń wykonano dla wąskiego zakresu zmienności parametrów CHM12, CHM22, C1RC2 oraz C2RC2. Wartości parametrów CHM12, CHM22 oraz C2RC2 losowano z rozkładu jednorodnego w przedziale 0.75 – 1.25. Natomiast wartości parametru C1RC2 próbkowano z przedziału 1.5 – 2.5 również z rozkładu jednorodnego. Pierwsza część obliczeń dla podstawowego zbioru uczącego przedstawiona jest na rysunku 5.6 kolorem szarym. Celem obliczeń dla pierwszej części zbioru było uzyskanie wyników zbliżonych do przebiegu eksperymentalnego, tak aby były one możliwe do przypisania do klas A, B oraz C.

Druga część obliczeń dla podstawowego zbioru uczącego wykonana została dla wartości parametrów CHM12, CHM22, C2RC2 z zakresu 0.05 – 2.0 oraz 0.05 – 4.0 dla C1RC2. Wartości losowane były z rozkładów jednorodnych. Wykonano 210 obliczeń, oznaczonych na rysunku 5.6 kolorem oliwkowym. Większość obliczeń z tego zbioru została zaklasyfikowana

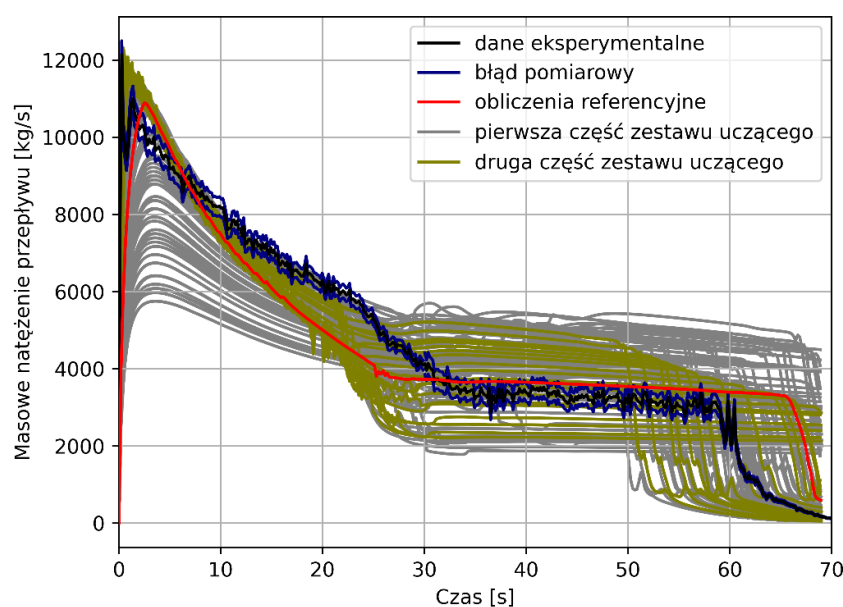
do klas D i E. Celem obliczeń dla tej części zestawu podstawowego było nauczenie algorytmu jak wyglądają przebiegi znacznie odbiegające od danych eksperymentalnych.

Test 8 Zestaw uczący



Rys. 5.6. Podstawowy zestaw uczący dla obliczeń masowego natężenia przepływu w instalacji Marviken w podziale na pierwszą część zestawu (kolor szary), oraz drugą część zestawu (kolor oliwkowy)

Test 27 Zestaw uczący



Rys. 5.7. Rozszerzony zestaw uczący dla obliczeń masowego natężenia przepływu w instalacji Marviken w podziale na pierwszą część zestawu (kolor szary), oraz drugą część zestawu (kolor oliwkowy)

Aby nie przetrenować algorytmu na tylko jednym zestawie danych, rozszerzono zbiór uczący o obliczenia dla drugiego testu. W tym celu wybrano test numer 27, dla którego wykonano łącznie 131 obliczeń, ponownie w dwóch częściach, co przedstawiono na rysunku 5.7. Pierwsze 104 obliczenia (kolor szary na rysunku 5.7) wykonano dla wartości parametrów losowanych z rozkładów jednorodnych o granicach 0.5 – 1.5 dla CHM12, CHM22, C2RC2 oraz 1.0 – 3.0 dla C1RC2. W związku z zauważalnym niedoszacowaniem wydatku przepływu w fazie przepływu cieczy przechłodzonej, zdecydowano o wykonaniu dodatkowych 27 obliczeń, gdzie parametr CHM12, mnożnik w fazie wypływu przechłodzonego, był losowany z rozkładu 1.5 – 2.5. Pozostałe parametry były losowane z poprzednio stosowanych rozkładów.

Każde z obliczeń ze zbioru podstawowego oraz rozszerzonego zostało porównane z danymi eksperymentalnymi i na tej podstawie manualnie dokonano przypisania do jednej z pięciu zdefiniowanych powyżej klas. Zgodnie z opisem w rozdziale 4.5.2.3 na tym etapie nie są wykorzystywane cechy, dokonywana jest więc jakościowa analiza wyników. Wyniki przypisania obliczeń do klas przedstawiono w tabeli 5.2. Dla zestawu podstawowego (test 8) 5.5% obliczeń zostało zaklasyfikowanych do klasy A. W zestawie rozszerzonym (test 8 oraz 27) udział obliczeń przypisanych do klasy A wzrósł do 7.3%. Najwięcej obliczeń przypisano do klasy D, czyli obliczeń, które nie są zgodne z danymi eksperymentalnymi zarówno w trakcie fazy wypływu przechłodzonego jak i dwufazowego.

Tabela 5.2. Liczba obliczeń przypisanych manualnie do danej klasy dla testów stanowiących podstawowy i rozszerzony zbiór uczący dla instalacji Marviken

Zestaw testowy	Klasa A	Klasa B	Klasa C	Klasa D	Klasa E	Suma
Test 8	40	135	95	388	62	720
Test 27	22	30	30	29	20	131
Suma	62	165	125	417	82	851

Następnie dla wszystkich 851 obliczeń z zestawu uczącego dokonano obliczeń cech opisanych powyżej, a w dalszej kolejności stosując klasyfikator lasów losowych (opisany w [rozdziale 4.5.2.4](#)) algorytm uczenia maszyn nauczył się relacji między cechami a klasami. Relacja ta będzie wykorzystywana w przypisywaniu klas przez algorytm dla obliczeń z zestawu testowego.

Zestaw testowy dla obliczeń masowego natężenia przepływu – instalacja testowa Marviken

Obliczenia wykonane dla czterech testów stanowiąc będą zestaw testowy. Testy różniły się warunkami początkowymi oraz kształtem dyszy, uzasadnienie wyboru testów przedstawiono przy opisie tabeli 5.1. Dla każdego z wybranych testów wykonano 60 000 obliczeń z zastosowaniem kodu TRACE. Wartości parametrów CHM12 oraz CHM22 były losowane

z rozkładu jednorodnego w przedziale 0.5 – 2.0. Wartości parametru C2RC2 były losowane z rozkładu jednorodnego w zakresie 0.01 – 2.5, natomiast parametru C1RC2 w zakresie 0.01 – 5.0. Wartości parametrów były określone poprzez próbkowanie z podanych rozkładów stosując generator wartości losowych Saltelliego. Rozkłady jednorodne stosowane są, aby podkreślić brak wiedzy na temat rzeczywistej wartości parametrów.

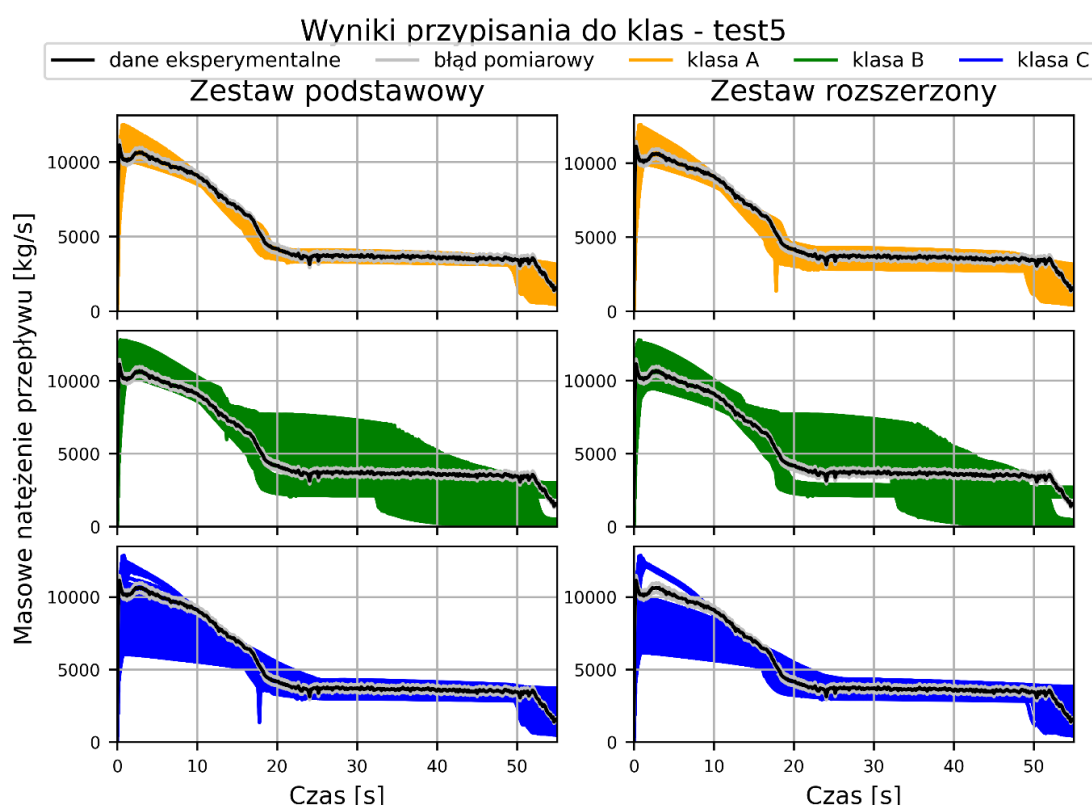
Po wykonaniu obliczeń, dla każdego z nich wyznaczone zostały wartości cech. Model uczenia maszyn znał również relację między cechami a klasami, której nauczył się podczas trenowania na zestawie uczącym. Na tej podstawie algorytm przypisywał jedną z pięciu klas dla każdego z wykonanych obliczeń, stosując klasyfikator lasów losowych. Trenowania algorytmu dokonano w pierwszej kolejności na podstawowym zestawie uczącym, a następnie na rozszerzonym zestawie uczącym, więc wyniki przypisania klas określono dla obu tych zestawów. W tabeli 5.3 przedstawiono liczbę obliczeń przypisanych przez model uczenia maszyn dla każdego z testów wpływu krytycznego w instalacji Marviken tworzących zbiór testowy, w podziale na dwa stosowane zestawy uczące.

Tabela 5.3. Liczba obliczeń przypisanych przez algorytm uczenia maszyn do danej klasy dla testów stanowiących zestaw testowy, stosując podstawowy i rozszerzony zbiór uczący dla instalacji Marviken

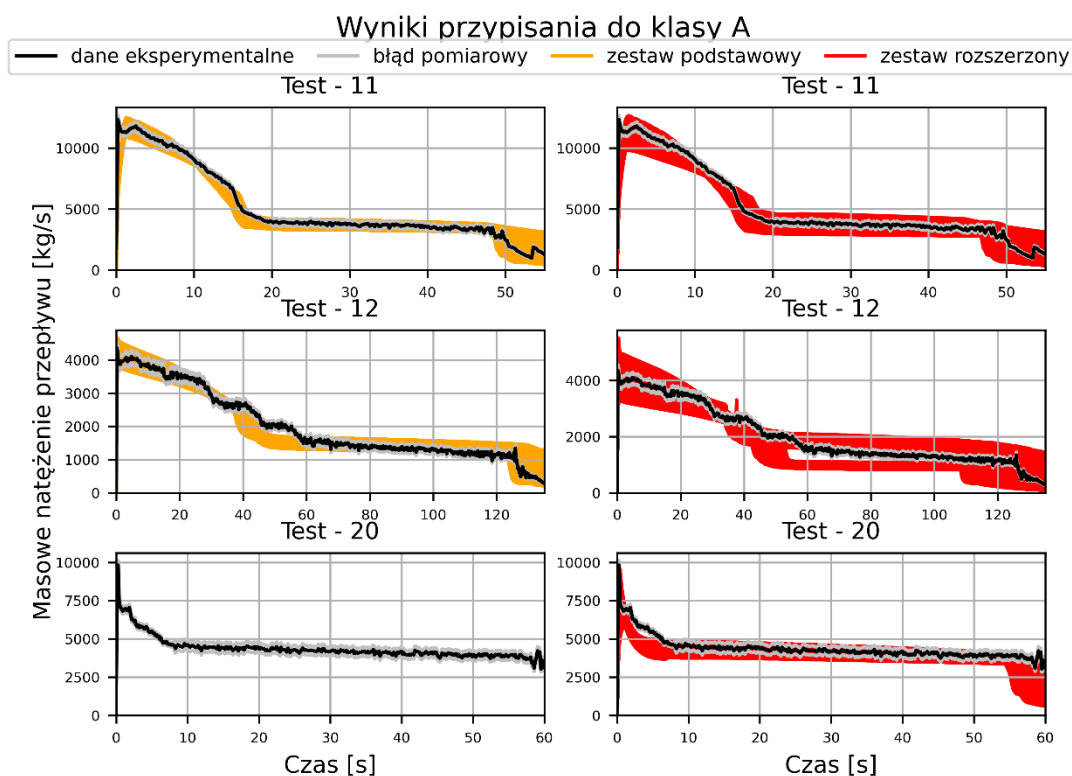
	Podstawowy zestaw uczący					Rozszerzony zestaw uczący				
	Klasa A	Klasa B	Klasa C	Klasa D	Klasa E	Klasa A	Klasa B	Klasa C	Klasa D	Klasa E
Test 5	2449	29807	5187	15618	6939	6893	22761	3856	21773	4717
Test 11	662	5693	7024	39557	7064	1357	5275	6932	40836	5600
Test 12	509	3376	13336	26000	16779	5107	3768	17949	33176	0
Test 20	0	14679	5817	27618	11886	3998	19810	3555	29237	3400
Suma	3620	53555	31364	108793	42668	17355	51614	32292	125022	13717
Udział procentowy	1.5%	22.3%	13.1%	45.3%	17.8%	7.2%	21.5%	13.5%	52.1%	5.7%

Stosując rozszerzony zestaw uczący do trenowania modelu uczenia maszyn osiągnięto większy procentowy udział obliczeń przypisanych do klasy A. Jest to szczególnie widoczne dla testu 20 dla którego stosując podstawowy zestaw uczący algorytm nie przypisał żadnemu z obliczeń klasy A. Dla wszystkich czterech testów z zestawu podstawowego przypisanie do klasy A było zdecydowanie bardziej restrykcyjne niż w przypadku zestawu rozszerzonego. Na rysunku 5.8 przedstawiono przypisanie do klas A, B, C wykonane z zastosowaniem modelu uczenia maszyn stosując podstawowy zestaw uczący (po lewej stronie na rysunku) oraz rozszerzony zestaw uczący (po prawej stronie na rysunku) dla testu wpływu krytycznego numer 5. Mimo, że dla rozszerzonego zestawu uczącego przypisano ponad 2 razy więcej obliczeń do klasy A, to zakres zmienności rezultatów w obu przypadkach jest bardzo podobny. Zakresy pokrywają cały przebieg zmian masowego natężenia przepływu w czasie, jednocześnie zakresy nie są zbyt

szerokie. W przypadku klasy B dla obu zestawów uczących można zauważyć, że algorytm poprawnie przypisał obliczenia, które są zbieżne z eksperymentem w fazie wypływu cieczy przechłodzonej. Natomiast w przypadku klasy C zbieżność jest odpowiednia tylko w fazie przepływu dwufazowego. Pokazuje to, że algorytm poprawnie przypisywał przebiegi do klas. Dla kolejnych testów, na rysunku 5.9 przedstawiono już tylko wyniki przypisania do klasy A, na podstawie których będą tworzone rozkłady gęstości prawdopodobieństwa parametrów modelowych. W przypadku testu 11 dla zestawu rozszerzonego przyporządkowano dwa razy więcej obliczeń, jednak w jakościowej analizie przebiegów nie można zaobserwować zauważalnej różnicy. Zakresy zmienności rezultatów pokrywają cały przebieg wydatku przepływu. W przypadku testu 12 do klasy A przypisano dziesięć razy więcej obliczeń, co zauważalnie wpływa na jakość dopasowania do danych eksperymentalnych. Szczególnie w obszarze przepływu dwufazowego obliczenia z zestawu rozszerzonego tworzą szeroki przedział zmienności, co nie jest obserwowalne w przypadku zestawu podstawowego. Dla testu 20 jedynie stosując zestaw rozszerzony udało się uzyskać dopasowanie prawie 4000 obliczeń do klasy A. Wyniki przypisania do wszystkich klas znajdują się w [aneksie](#).



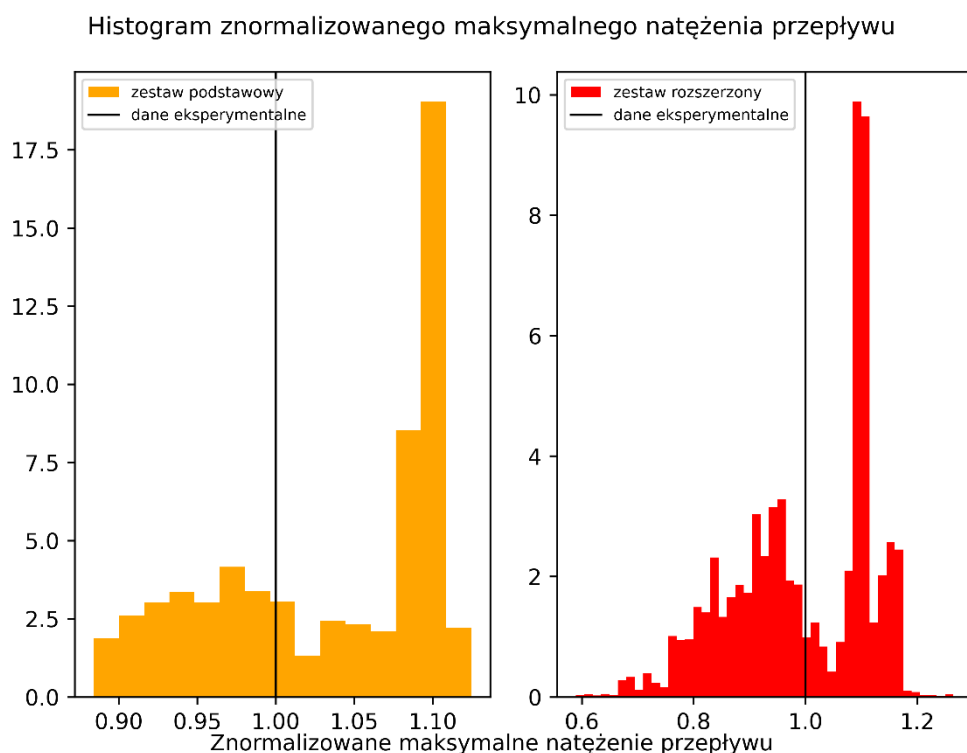
Rys. 5.8. Przypisanie do klas A, B, C wykonane z zastosowaniem modelu uczenia maszyn stosując podstawowy zestaw uczący oraz rozszerzony zestaw uczący dla testu 5 w instalacji Marviken



Rys. 5.9. Przypisanie do klas A dla testów 11, 12 oraz 20 wykonane z zastosowaniem modelu uczenia maszyn stosując podstawowy zestaw uczący oraz rozszerzony zestaw uczący w instalacji Marviken

Ogólnie w jednym przypadku zastosowanie zestawu podstawowego do nauki algorytmu uczenia maszyn spowodowało zdecydowanie dokładniejsze przypisanie do klasy A (test 12), w dwóch przypadkach różnice nie były aż tak znaczące (testy 5 i 11), natomiast dla testu 20 jedynie stosując zestaw rozszerzony można było uzyskać przypisanie do klasy A. Pokazuje to wrażliwość algorytmu na zestaw uczący, oraz większą uniwersalność rozszerzonego zestawu uczącego składającego się z obliczeń dla różnych testów. Uniwersalność ta z drugiej strony osiągana jest częściowo kosztem restrykcyjności oraz dokładności dopasowania do danych eksperymentalnych.

Podobnie jak dla obliczeń wykonywanych stosując metodę Bayesowską wyznaczono wartości znormalizowanego maksymalnego natężenia przepływu dla wszystkich obliczeń przypisanych do klasy A, co przedstawiono na rysunku 5.10. Oszacowane maksymalne natężenie przepływu było w większości przypadków przeszacowane o 10 do 15 procent. Osiągnięto więc lepsze dopasowanie do danych eksperymentalnych niż w przypadku stosowania metody opartej o analizę Bayesowską. Wynika to z faktu, że jedną z cech jaką algorytm ewaluował przy przypisywaniu klas do wyników obliczeń jest maksymalne natężenie przepływu. Dodatkowo algorytm w wielu przypadkach oceniał tą cechę jako najistotniejszą w przypisywaniu obliczeń do klas.

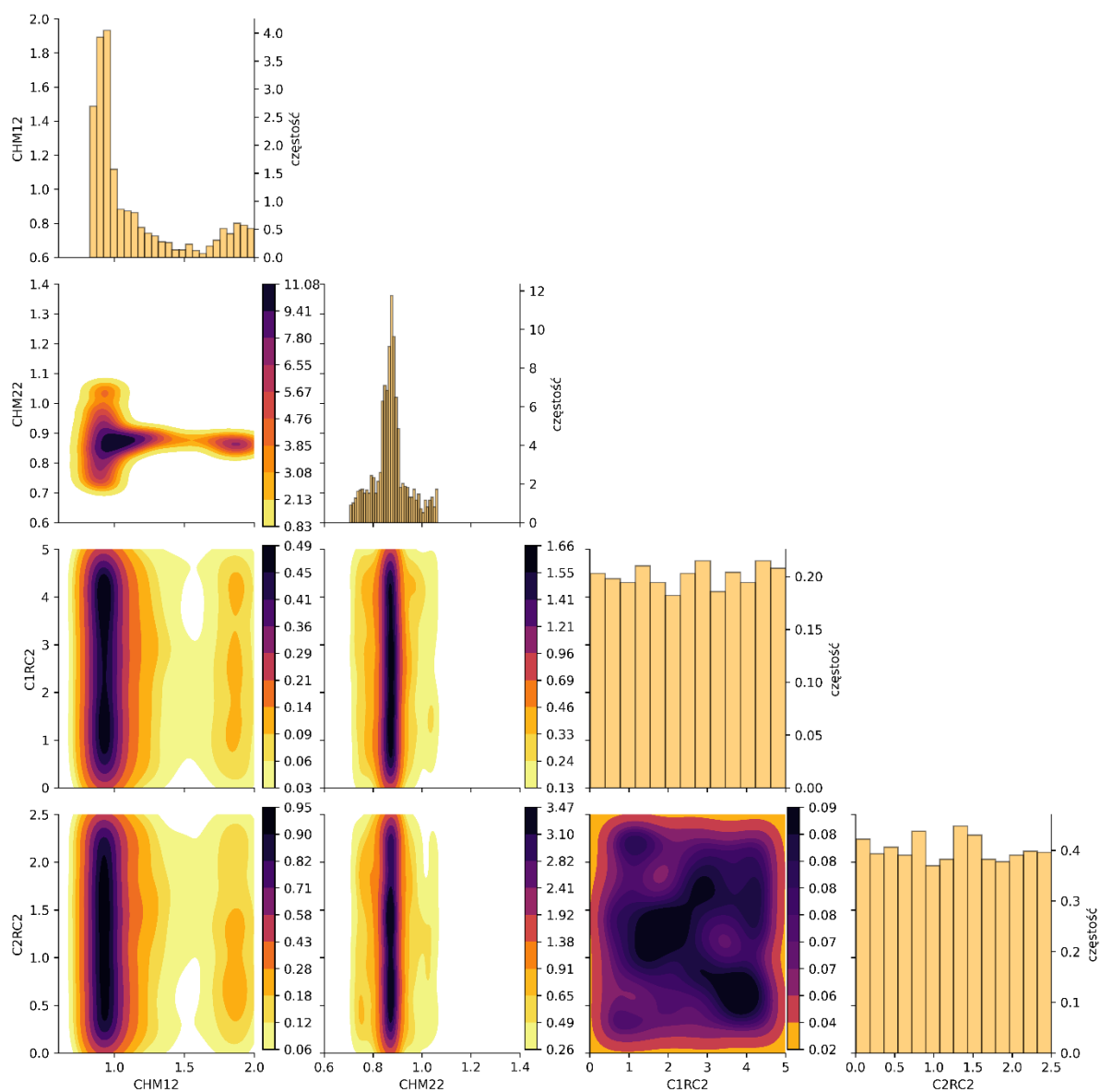


Rys. 5.10. Znormalizowany wydatek przepływu dla obliczeń zaklasyfikowanych do klasy A dla algorytmu uczenia maszyn dla testów 4, 11, 12 i 20 w instalacji Marviken

Określenie rozkładów gęstości prawdopodobieństwa parametrów

Wszystkie obliczenia wykonywane dla zbioru testowego wykonywane były na podstawie plików wsadowych w których określone były wartości parametrów modelowych CHM12, CHM22, C1RC2 oraz C2RC2. Jedynie te wartości, losowane z rozkładów omówionych w poprzednim podrozdziale, stanowiły różnice w plikach wsadowych dla wszystkich obliczeń. Dla wszystkich testów obliczenia wykonywano z tymi samymi zestawami parametrów. Tak więc, wartości parametrów, dla których wykonane obliczenia zostały przypisane do klasy A stanowią rozwiązanie problemu wstecznej kwantyfikacji niepewności. Należało jednak odrzucić uprzednio wszystkie obliczenia dla których przypisano choć jedną klasę D. Tak ograniczone zbiory wartości każdego z parametrów dla obliczeń zaklasyfikowanych pozwoliły na stworzenie rozkładów gęstości prawdopodobieństwa tych parametrów. Rozkłady gęstości prawdopodobieństwa dla parametrów CHM12, CHM22, C1RC2 oraz C2RC2 uzyskane wykorzystując do uczenia algorytmu podstawowy zbiór testowy przedstawiono na rysunku 5.11, natomiast stosując rozszerzony zbiór testowy na rysunku 5.12. Podobnie jak w przypadku pierwszej metody kwantyfikacji niepewności, przedstawiono również wykresy dwóch zmiennych, czyli dwuwymiarowe rozkłady pokazujące zakresy wartości w których

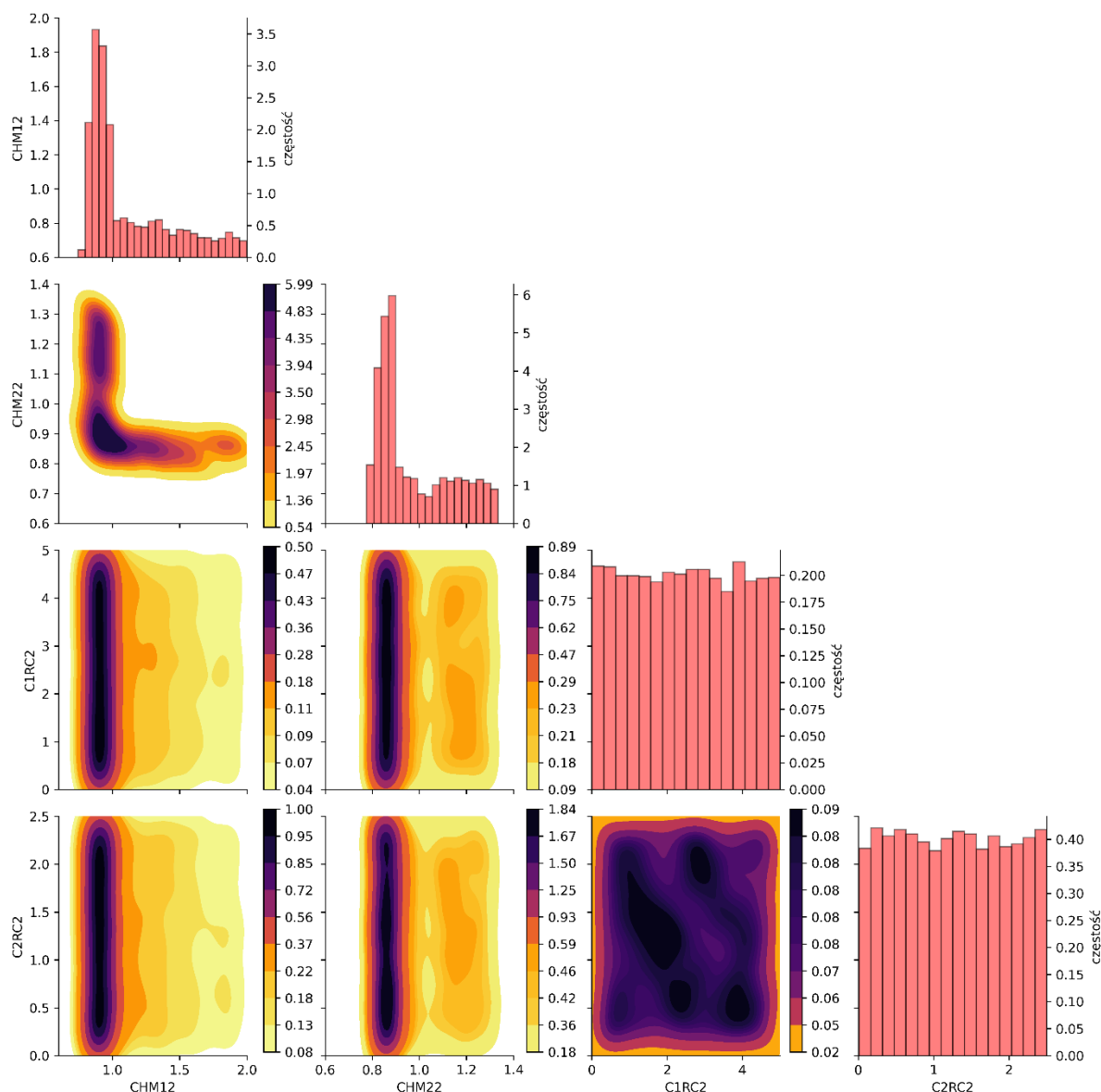
prawdopodobieństwo uzyskania wyników obliczeń zbliżonych do danych eksperymentalnych jest największe.



Rys. 5.11. Jedno i dwuwymiarowe rozkłady gęstości prawdopodobieństwa dla parametrów modelowych CHM12, CHM22, C1RC2, C2RC2 uzyskane stosując metodę wstecznej kwantyfikacji niepewności opartą o algorytm uczenia maszyn trenowany na podstawowym zestawie uczącym

Dla obu zestawów uczących wartości parametru CHM12 w zakresie pomiędzy 0.8 a 1.0 są najbardziej prawdopodobne. Częstości występowania pozostałych wartości w przedziale 0.7-2.0 są ponad cztery razy mniejsze. W obu przypadkach nie zarejestrowano wartości poniżej 0.7. W przypadku parametru CHM22 ponownie otrzymano wartości w ograniczonym przedziale – od 0.7 do 1.3 w przypadku zestawu rozszerzonego oraz od 0.7 do 1.1 w przypadku zestawu podstawowego. W obu przypadkach rozkład posiada charakterystyczne maksima w okolicach

wartości 0.86. W rozkładzie dla zestawu podstawowego zauważalne jest, że pozostałe wartości są ponad dziesięciokrotnie mniej prawdopodobne, w przypadku zestawu rozszerzonego pozostałe wartości są sześć razy mniej prawdopodobne a rozkład jest szerszy. Wynika to ze znacznej różnicy w liczbie obliczeń przyporządkowanych do klasy A. Stosując zestaw rozszerzony przypisano do klasy A 17355 obliczeń, zestaw podstawowy tylko 3620. W obu przypadkach rozkłady są jednak bardzo wąskie, wartości powyżej 1.3 nie powinny być więc stosowane.



Rys. 5.12. Jedno i dwuwymiarowe rozkłady gęstości prawdopodobieństwa dla parametrów modelowych CHM12, CHM22, C1RC2, C2RC2 uzyskane stosując metodę wstecznej kwantyfikacji niepewności opartą o algorytm uczenia maszyn trenowany na rozszerzonym zestawie uczącym

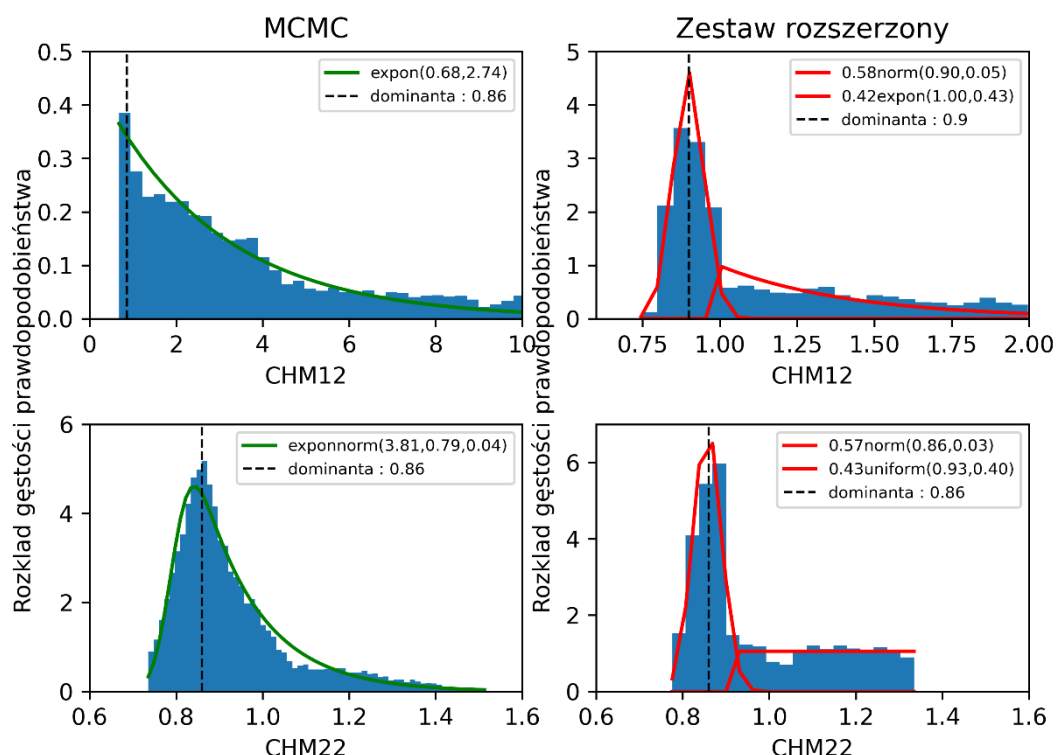
Niezależnie od zastosowanego zestawu uczącego rozkłady gęstości prawdopodobieństwa dla parametrów C1RC2 oraz C2RC2 posiadają rozkłady jednorodne w zakresie ich zmienności określonym a’priori. Rozkłady a’posteriori są więc identyczne jak rozkłady a’priori. Dokonując analizy ich rozkładów dla pozostałych klas, również uzyskiwano rozkłady jednorodne. Sugeruje to, że wartości tych parametrów nie miały wpływu na wyniki obliczeń, które porównywane były z danymi eksperymentalnymi. Na podstawie wykonanej analizy nie można więc określić wartości najlepszego szacowania dla pary tych parametrów, ponieważ w badanym zakresie każda z wartości jest tak samo prawdopodobna.

5.1.4.3. Porównanie wyników oraz propozycja rozkładów gęstości prawdopodobieństwa

Do końcowej analizy rozkładów gęstości prawdopodobieństwa wybrano ostatecznie obliczenia uzyskane po zastosowaniu rozszerzonego zestawu uczącego. Pozwolił on na przypisanie większej liczby obliczeń do klasy A, w tym dla testu 20 gdzie dla zestawu podstawowego nie przypisano żadnych obliczeń. Dodatkowo został on utworzony na podstawie obliczeń dla dwóch różnych testów w zbiorze uczącym, nie jest więc przetrenowany na danych dla jednego testu o specyficznych warunkach początkowych i brzegowych. W ten sposób rezultaty przypisania do klas są bardziej dostosowane do zróżnicowanych przebiegów wartości wielkości wyjściowej.

Aby zaprezentowane na rysunkach 5.5 oraz 5.12, uzyskane z zastosowaniem dwóch metod kwantyfikacji niepewności, rozkłady gęstości prawdopodobieństwa dla parametrów modelowych mogły być zastosowane w przyszłych obliczeniach konieczne jest opisanie ich funkcjami ciągłymi. Funkcje te będą reprezentacją dopasowanych propozycji rozkładów. W tym celu wykorzystano pakiet SciPy [102] dostępny w języku programowania python. W pakiecie tym dostępna jest funkcja `rv_continuous.fit` opierająca się o estymację metodą największej wiarygodności. Funkcja ta umożliwia znajdowanie wartości parametrów dla zaproponowanego rozkładu, tak aby rozkład ten w jak najlepszy sposób odzwierciedlał dane wejściowe. Danymi wejściowymi są w tym wypadku wszystkie wartości parametru tworzące histogramy przedstawione na rysunkach 5.5 oraz 5.12. Następnie stosując funkcję `rv_continuous.fit` porównywano wiele rozkładów poszukując takiego który w najlepszym stopniu odwzoruje dane wejściowe. Weryfikacji poprawności dopasowania rozkładu dokonywano przez porównanie dystrybuant. Ostatecznie dopasowane rozkłady dla parametrów modelowych CHM12 oraz CHM22 przedstawiono na rysunku 5.13 w porównaniu z histogramami rzeczywistych wartości parametrów. W przypadku metody Bayesowskiej zaproponowano rozkłady wykładniczy oraz wykładniczo-normalny. Dla obu parametrów

dominanta wynosiła 0.86. Dla metody uczenia maszyn zastosowano złożenie dwóch rozkładów dla obu parametrów. Dla CHM12 złożono rozkład normalny oraz wykładniczy, ponieważ można zaobserwować, że wartości mogą wybiegać poza założoną a priori górną granicę zakresu równą 2. W przypadku CHM22 zastosowano rozkład normalny i jednorodny zakończony na 1.33, gdyż algorytm nie wskazał, że prawdopodobne są jakiekolwiek wartości, poza tym zakresem.

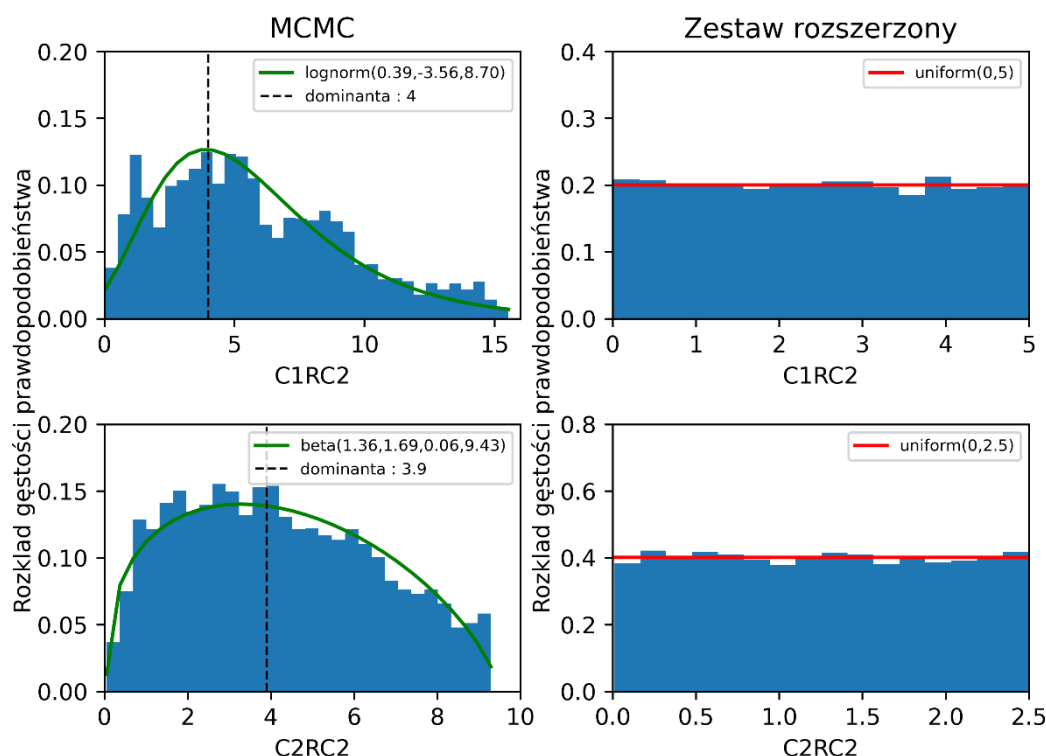


Rys. 5.13. Propozycje ciągłych rozkładów gęstości prawdopodobieństwa dla parametrów CHM12 oraz CHM22 w wyniku zastosowania dwóch metod wstecznej kwantyfikacji niepewności

Dla parametrów C1RC2 oraz C2RC2 określono rozkłady wykładniczo-normalny oraz logarytmiczno-normalny dla metody opartej o analizę Bayesowską, co przedstawiono na rysunku 5.14. Dominanta dla parametru C1RC2 wynosiła 1.25, natomiast dla parametru C2RC2 3.9. W przypadku metody opartej o uczenie maszyn rozkłady dla obu parametrów są jednorodne. Dla rozkładów takich nie można określić dominanty, stąd uznano, że wartością oczekiwaną jest wartość domyślna, czyli 2.0 dla C1RC2 oraz 1.0 dla C2RC2.

Przedstawione rozkłady gęstości prawdopodobieństwa parametrów będą stanowiły podstawę do losowaniach z nich wartości parametrów w obliczeniach sprawdzających i walidacyjnych. Bardziej szczegółowe omówienie porównawcze uzyskanych wyników i rozkładów znajduje się

w artykule opublikowanym na podstawie wyników badań przeprowadzonych w ramach rozprawy [103].



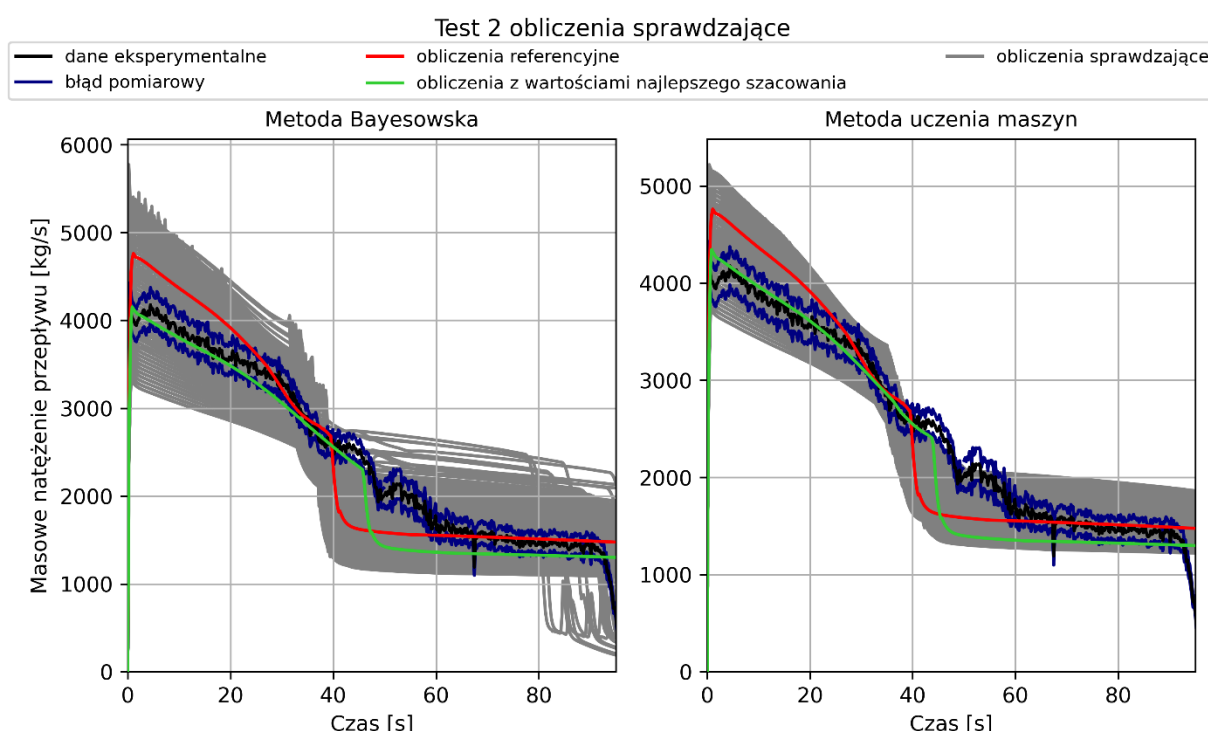
Rys. 5.14. Propozycje ciągłych rozkładów gęstości prawdopodobieństwa dla parametrów C1RC2 oraz C2RC2 w wyniku zastosowania dwóch metod wstecznej kwantyfikacji niepewności

5.1.5. Sprawdzenie poprawności uzyskanych rozkładów gęstości prawdopodobieństwa (krok 13)

Krok trzynasty metodyki to pierwsze sprawdzenie uzyskanych rozkładów gęstości prawdopodobieństwa dla parametrów modelowych. W tym celu zwyczajowo wykorzystuje się testy które nie zostały wykorzystane w procesie uzyskiwania rozkładów. W przypadku instalacji Marviken, gdzie czas obliczeniowy jest krótki, wykonano obliczenia sprawdzające dla wszystkich 26 dostępnych testów stosując dwa zestawy rozkładów – dla metody opartej o analizę Bayesowską oraz dla metody opartej o uczenie maszyn. W obu wypadkach wykonano po 1000 obliczeń dla każdego testu, a wartości parametrów były losowane z rozkładów opisanych w poprzedniej sekcji. Należy zaznaczyć, że obliczenia wykonywane w tym kroku nie są pełną analizą niepewności, ponieważ jedynymi zmienianymi wartościami parametrów będą parametry modelowe CHM12, CHM22, C1RC2 oraz C2RC2. Pozwalają one jednak na wstępną ocenę poprawności wyznaczonych rozkładów niepewności wymienionych parametrów. Losowania parametrów dokonano stosując proste próbkowanie losowe, a przedstawione rezultaty przedstawiają 95 percentyl zgodnie z formułą Wilksa omówioną

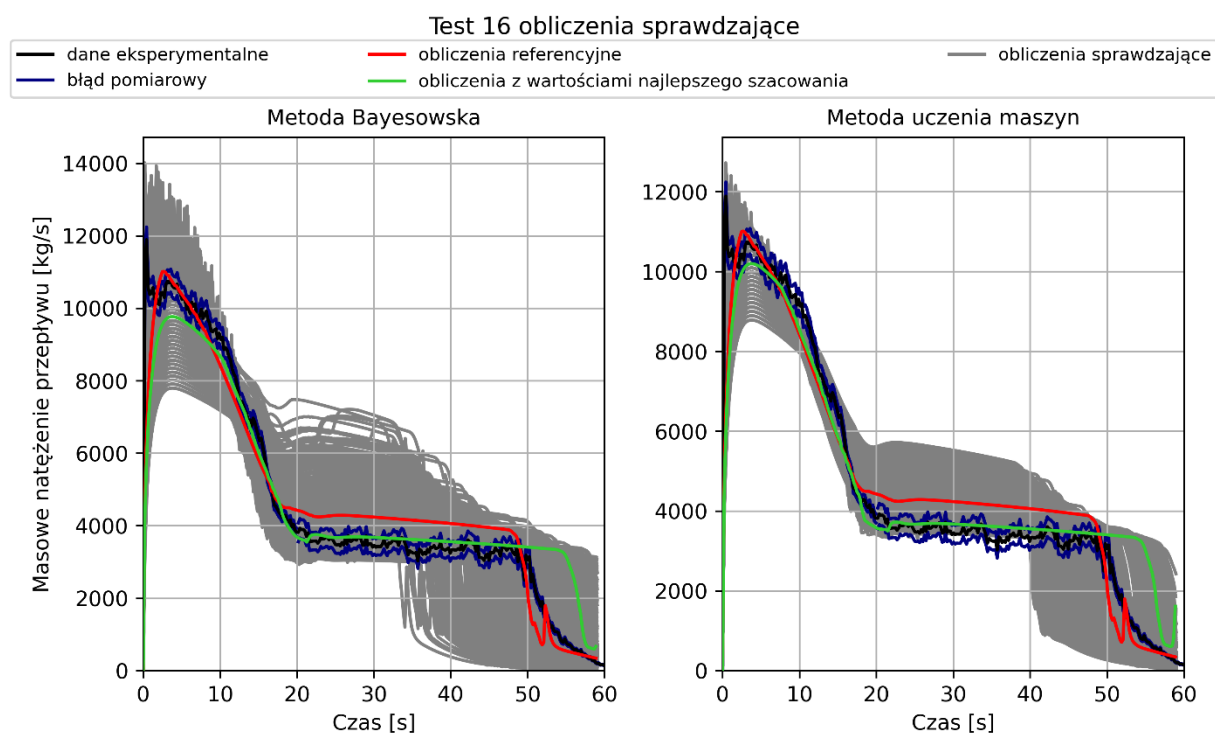
w [rozdziale 2.1.1](#). Dodatkowo dla każdego testu wykonano obliczenia z wartościami najlepszego szacowania parametrów – czyli dominantami dla rozkładów przedstawionych w poprzednim podrozdziale.

Wyniki obliczeń sprawdzających przedstawiono w podziale na grupy zdefiniowane w tabeli 5.1. Dla każdej z grup wybrano jeden reprezentatywny przypadek. Na rysunku 5.15 przedstawiono wyniki obliczeń sprawdzających dla testu numer 2, wybranego jako reprezentatywny z grupy pierwszej. Obliczenia referencyjne dla grupy pierwszej charakteryzowały się błędnym wyznaczaniem czasu zmiany charakteru przepływu z jednofazowego na dwufazowy. Obliczenia z wartościami najlepszego szacowania parametrów modelowych pozwoliły na bardziej poprawne odzwierciedlenie czasu zmiany charakterystyki przepływu a także na lepsze oszacowanie przebiegu w fazie przepływu przechłodzonego. Zakresy niepewności, powstałe z nałożenia wszystkich otrzymanych przebiegów z obliczeń sprawdzających, dla obu zastosowanych metod obejmują dane eksperymentalne wraz z błędem pomiarowym. Uzyskane zakresy niepewności są szersze dla metody Bayesowskiej dla wszystkich obliczeń w tej grupie, w wielu wypadkach zdecydowanie przeszacowując dane eksperymentalne, co można zaobserwować na rysunku 5.15.



Rys. 5.15. Zakres niepewności przebiegu masowego natężenia przepływu utworzony przez wyniki obliczeń sprawdzających dla reprezentatywnego testu z grupy pierwszej

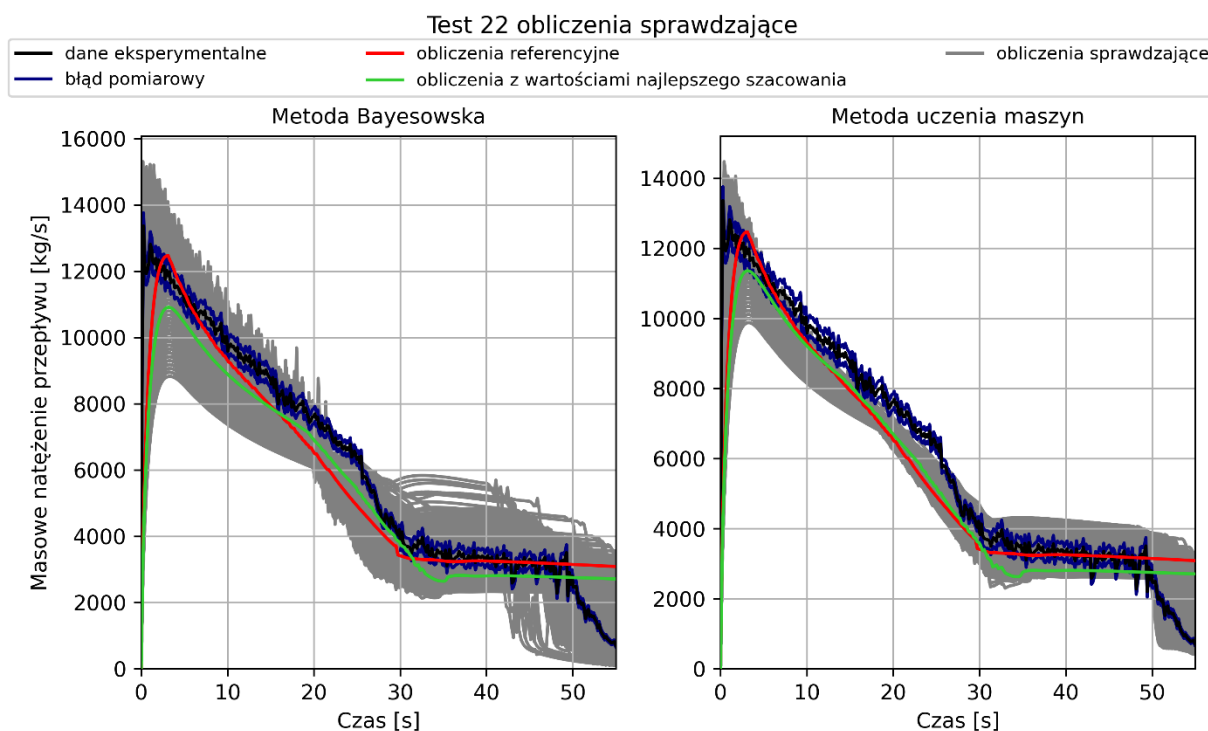
Na rysunku 5.16 przedstawiono wyniki obliczeń dla testu 16 będącego reprezentatywnym dla grupy drugiej. Spośród grupy drugiej pochodziły 2 testy wykorzystane w zestawie testowym, jeden w zestawie uczącym oraz dwa wykorzystane w obliczeniach MCMC, stąd oczekiwano, że dla tej grupy wyniki obliczeń z parametrami najlepszego szacowania będą najbliższe danym eksperymentalnym, a zakresy nie będą nadmiernie szerokie. Obliczenia z wartościami najlepszego szacowania zdecydowanie lepiej odwzorowują przebieg eksperymentu niż obliczenia referencyjne. Jednakże uzyskane zakresy niepewności, w fazie przepływu dwufazowego są bardzo szerokie. Ponownie to obliczenia uzyskane stosując metodę Bayesowską dają w efekcie szersze zakresy niepewności. W fazie przepływu przechłodzonego zakres niepewności uzyskany metodą uczenia maszyn jest wąski i w odpowiedni sposób pokrywa dane eksperymentalne. Jest to więc druga grupa, dla której uzyskane zakresy są bardziej zachowawcze dla metody uczenia maszyn.



Rys. 5.16. Zakres niepewności przebiegu masowego natężenia przepływu utworzony przez wyniki obliczeń sprawdzających dla reprezentatywnego testu z grupy drugiej

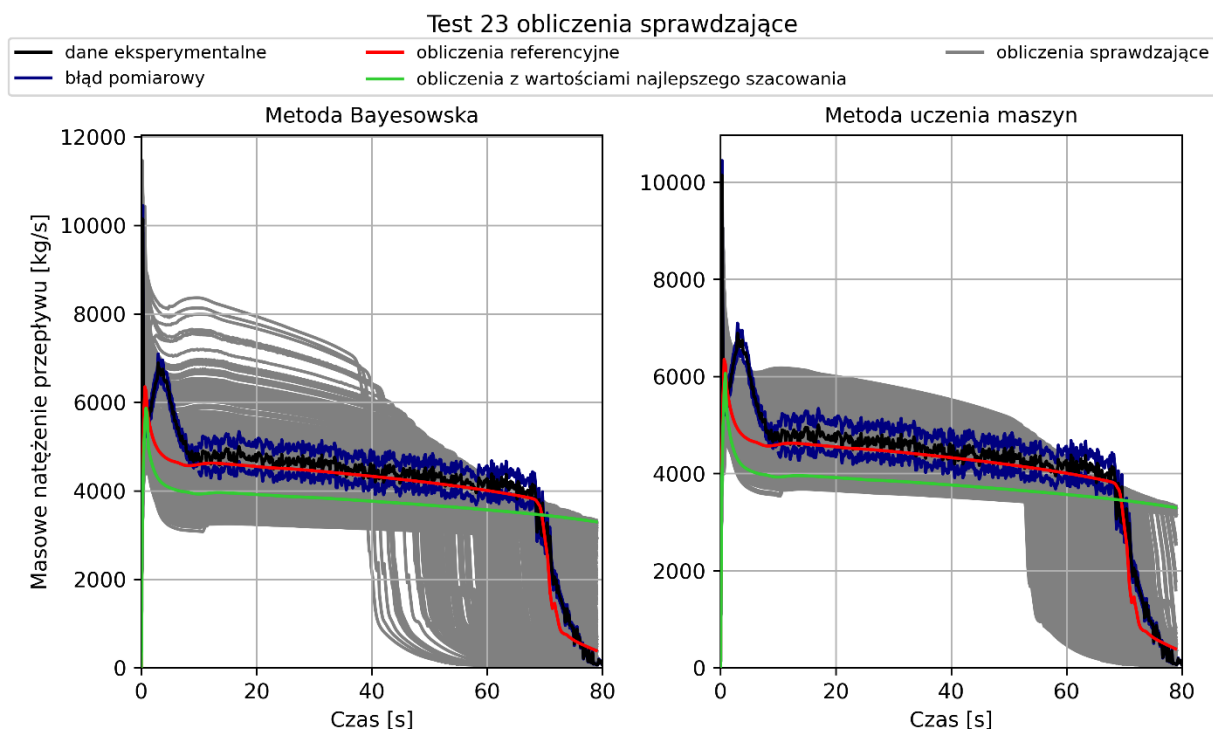
Na rysunku 5.17 zaprezentowano wyniki obliczeń sprawdzających dla testu 22 będącego reprezentatywnym dla grupy trzeciej. Spośród tej grupy tylko test nr 22 nie był wykorzystywany w zestawie uczącym, testowym bądź obliczeniach MCMC. Obliczenia z wartościami najlepszego szacowania pozwoliły na osiągnięcie poprawy w dopasowaniu do danych eksperymentalnych w porównaniu do obliczeń referencyjnych. Uzyskane zakresy

niepewności dla metody uczenia maszyn są wąskie a jednocześnie pokrywają dane eksperymentalne wraz z błędem pomiarowym. W przypadku metody Bayesowskiej ponownie zakresy są szerokie w fazie przepływu dwufazowego. Zakresy jednak w pełni obejmują dane eksperymentalne i błąd pomiarowy, zachowując prawidłowy kształt przebiegu.



Rys. 5.17. Zakres niepewności przebiegu masowego natężenia przepływu utworzony przez wyniki obliczeń sprawdzających dla reprezentatywnego testu z grupy trzeciej

Na rysunku 5.18 zaprezentowano wyniki obliczeń sprawdzających dla testu nr 23 wybranego z grupy czwartej. Grupa ta zawiera obliczenia, które przebiegiem odbiegają od innych grup. Dla większości przypadków zarówno obliczenia referencyjne jak i obliczenia z wartościami najlepszego szacowania nie dają wyników zbliżonych z eksperymentalnymi. Dane eksperymentalne pokrywane są przez wyniki obliczeń sprawdzających dla wszystkich testów poza testem nr 14. Jednakże dla wszystkich obliczeń osiągnięto szerokie zakresy niepewności, szczególnie dla testów nr 9, 22 oraz 25. Dla testów nr 14 i 25 osiągnięte zakresy nie były nadmiernie szerokie. Stąd pozytywne jest, że mimo, że zakresy są bardzo szerokie to ostatecznie pokrywają dane eksperymentalne.



Rys. 5.18. Zakres niepewności przebiegu masowego natężenia przepływu utworzony przez wyniki obliczeń sprawdzających dla reprezentatywnego testu z grupy trzeciej

W większości przypadków wyniki obliczeń sprawdzających pozwoliły na pokrycie zakresu błędu pomiarowego danych eksperymentalnych. Ze względu na wykorzystanie wielu testów należących do różnych grup podczas wyznaczania rozkładów, spodziewane było, że zakresy zmienności wyników będą szerokie. Wciąż, szczególnie dla grup 1-3 uzyskane przebiegi w sposób jakościowy poprawnie odwzorowują przebiegi masowego natężenia przepływu. Węższe zakresy uzyskano stosując metodę opartą o uczenie maszyn, natomiast w większości przypadków to wartości najlepszego szacowania uzyskane w metodzie Bayesowskiej pozwalały na uzyskanie wyników bardziej zbieżnych z danymi pomiarowymi.

5.2. Obliczenia dla instalacji FLECHT SEASET

Instalacja testowa FLECHT SEASET (Full-Length Emergency Core Heat Transfer - Separate Effects And System Effects Test) służyła jako źródło danych eksperymentalnych dla badania zjawisk związanych szczególnie z wymianą ciepła w rdzeniu reaktora podczas fazy zalewania w trakcie awarii typu LBLOCA [104]. Badano między innymi wymianę ciepła w warunkach naturalnej cyrkulacji, chłodzenia parą wodną przy niskich wartościach liczby Reynoldsa oraz w warunkach wymuszonego zalewania rdzenia. Testy przeprowadzono w warunkach odwzorowujących odkrycie rdzenia reaktora, więc nawet przy wartościach temperatury oprzyrządowania symulującego pręty paliwowe przekraczających 1200°C. Kampania

eksperymentalna w instalacji Flecht Seaset była między innymi następstwem awarii w Three Mile Island, gdy zaczęto poszukiwać możliwości zbadania zjawisk zachodzących w rdzeniu podczas jego odkrywania. W niniejszej rozprawie wykorzystano dane eksperymentalne dla wymuszonego zalewania rdzenia dla niezablokowanych kaset paliwowych.

5.2.1. *Opis modelowanych zjawisk*

Faza zalewania rdzenia w trakcie awarii typu LBLOCA rozpoczyna się w momencie, gdy dolna komora zbiornika reaktora (ang. „lower plenum”) zostaje wypełniona wodą wtryskiwaną przez systemy bezpieczeństwa i woda zaczyna zalewać rdzeń reaktora. W trakcie fazy zalewania rdzenia odsłonięte w poprzedniej fazie awarii kasety paliwowe zostają więc ponownie schładzane. Odsłonięcie prętów paliwowych, a następnie ich ponowne zalanie wodą powoduje gwałtowne zmiany w charakterze wymiany ciepła na ściankach prętów paliwowych. W trakcie fazy wydmuchu dochodzi do stagnacji przepływu w rdzeniu co prowadzi do chwilowego kryzysu wrzenia i znaczącego pogorszenia wymiany ciepła. Ustalają się mechanizmy wymiany ciepła takie jak wrzenie błonowe oraz wrzenie przejściowe. Wtrysk wody z systemów bezpieczeństwa powoduje wznowienie przepływu przez rdzeń. Zalanie dolnej komory oraz dolnej części rdzenia powoduje powstanie przepływu pary, która odbiera ciepło od koszulek paliwowych. Następnie obserwowana jest propagacja frontu zalewania aż do górnej części rdzenia. W zależności od warunków przepływu (prędkości oraz stopnia przechłodzenia wody) powyżej poziomu frontu zalewania ustala się odwrócony przepływ pierścieniowy (ang. „inverted annular flow”) bądź rozproszony przepływ kropelkowy (ang. „dispersed droplet flow”). Po zetknięciu się chłodnej wody z koszulkami paliwowymi wymiana ciepła następuje w wyniku wrzenia pęcherzykowego, dzięki czemu znacznie wzrasta współczynnik wymiany ciepła. Kiedy dochodzi do momentu, gdy temperatura koszulki paliwowej osiąga temperaturę nasycenia dla ciśnienia panującego w obudowie bezpieczeństwa (bądź innej komorze, do której następuje wypływ chłodziwa np. w eksperymencie) to oznacza, że rdzeń będzie mógł być długoterminowo schładzany poprzez konwekcję.

Powyższy opis jest jedynie generalizacją i uproszczeniem wszelkich zjawisk zachodzących w rdzeniu podczas awarii typu LBLOCA, pozwala jednak zwrócić uwagę na wielość mechanizmów wymiany ciepła oraz warunków przepływu zachodzących w fazie zalewania. Dwoma wartościami, które w mierzalny sposób określają wymianę ciepła w rdzeniu w fazie zalewania są temperatura koszulki paliwowej oraz propagacja frontu zalewania (poziom wody w rdzeniu). Temperatura koszulki paliwowej a szczególnie jej maksymalna wartość w całym czasie jest jedną z najważniejszych wartości stanowiących kryterium akceptacji dla awarii typu

LBLOCA. Według wymagań amerykańskiej komisji dozoru jądrowego maksymalna temperatura koszulki paliwowej nie powinna przekroczyć 1204 °C. Przebieg zmian wartości temperatury koszulki paliwowej pozwala również na obserwację momentów czasowych w których dochodzi do zmian charakteru wymiany ciepła. Propagacja frontu zalewania prezentowana jest w formie przebiegu poziomu wody w rdzeniu w czasie awarii.

W kodzie TRACE modele wymiany ciepła przez ściankę podzielone są na trzy główne grupy powiązane z wystąpieniem kryzysu wrzenia [51]:

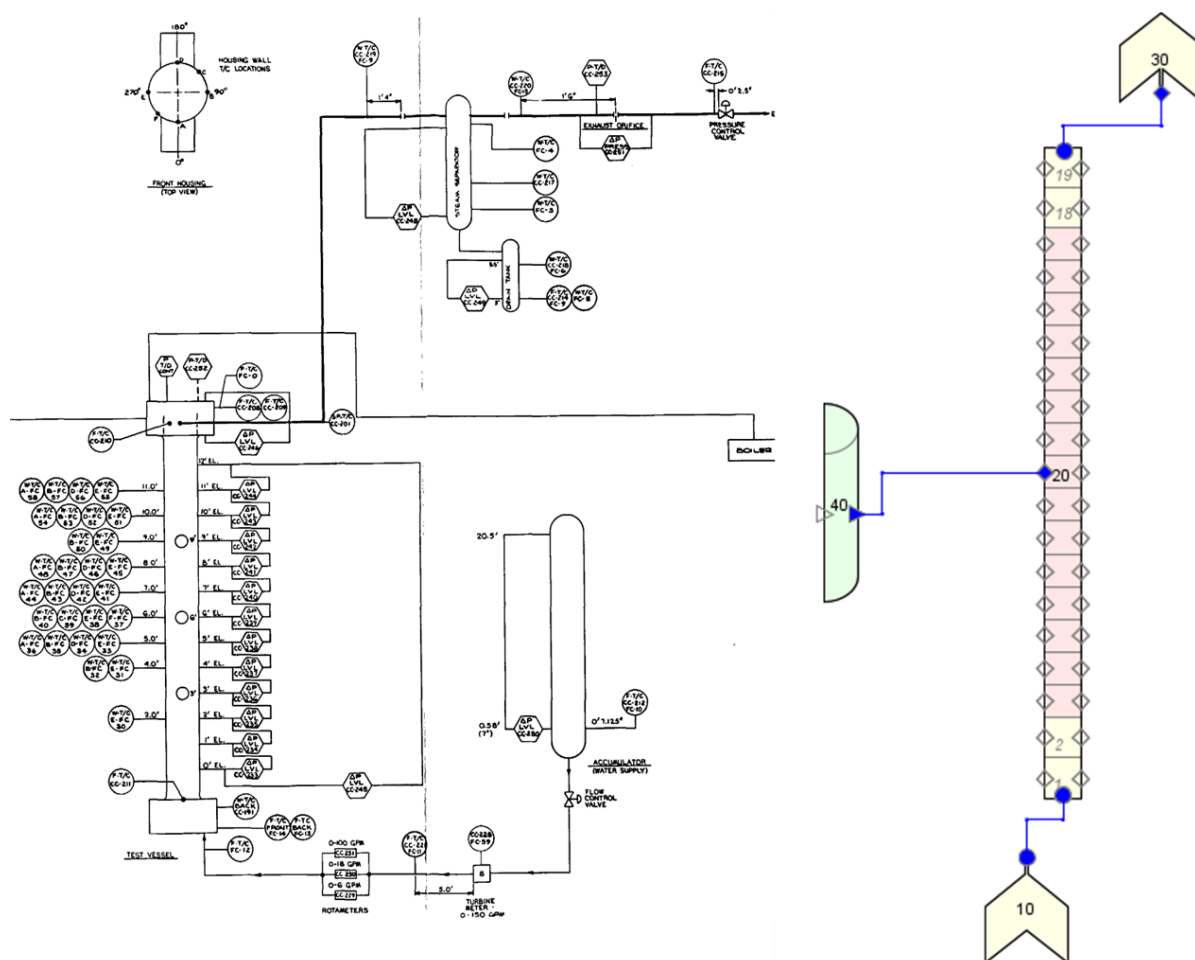
- Wymiana ciepła przed kryzysem wrzenia – obejmuje konwekcję w warunkach przepływu jednofazowego i dwufazowego, wrzenie pęcherzykowe, początek kryzysu wrzenia oraz wrzenie przechłodzone;
- Kryzys wrzenia – wyznaczenie krytycznego strumienia cieplnego;
- Wymiana ciepła po kryzysie wrzenia – wrzenie przejściowe, wrzenie błonowe (przy odwróconym przepływie pierścieniowym, rozproszonym przepływie kropelkowym, odwróconym przepływie korkowym).

TRACE pozwala na śledzenie zmian charakteru przepływu oraz mechanizmów wymiany ciepła, jednak w celu obserwacji zmian warunków panujących w rdzeniu wygodniejsze są parametry takie jak propagacja frontu zalewania czy temperatura koszulki paliwowej. Propagacja frontu zalewania jest parametrem wyjściowym dla struktur cieplnych stosowanych w kodzie TRACE do modelowania wymiany ciepła.

5.2.2. Opis instalacji eksperymentalnej FLECHT SEASET, kwalifikacja modelu (krok 8)

Główną częścią instalacji eksperymentalnej Flecht Seaset jest pęczek grzałek symulujących kasetę paliwową. Pęczek składał się z 161 prętów oraz szesnastu rurek prowadzących. Układ prętów oraz rurek prowadzących odzwierciedlał konfigurację zestawu paliwowego firmy Westinghouse typu 17x17. Pęczek grzałek znajdował się w zbiorniku o średnicy 0.194 metra. Część podlegająca grzaniu miała długość 3.66 metra, grzałki były elektryczne poprzez umieszczenie nagrzewnic kantalowych w prętach o średnicy 9.5 mm. Całkowita moc zapewniana do grzania prętów wynosiła 850 kW. Pęczek zawierał również osiem kratek dystansujących. Poniżej grzałek znajdowała się dolna komora, a powyżej górna komora. Na instalację eksperymentalną składają się również zbiornik i rurociągi wody chłodzącej, zbiornik końcowy oraz kocioł parowy. Testy warunków zalewania rdzenia przeprowadzano poprzez nagrzewanie prętów do odpowiedniej temperatury oraz dostarczając parę z kotła parowego, aby ustalić odpowiednie ciśnienie. Następnie w momencie rozpoczęcia testu woda ze zbiornika

wody chłodzącej trafiała do dolnej komory, a następnie zalewała pęczek grzałek. Przeprowadzono łącznie 63 udane testy zalewania zestawu paliwowego z różnymi warunkami początkowymi.



Rys. 5.19. Instalacja testowa Flecht Seaset [104] oraz model instalacji opracowany z zastosowaniem kodu TRACE w środowisku graficznym SNAP

Na rysunku 5.19 przedstawiono konfigurację instalacji testowej Flecht Seaset dla badania zjawiska zalewania rdzenia oraz opracowany w kodzie obliczeniowym TRACE i środowisku graficznym SNAP model instalacji testowej. Opracowany model posiada szczegółową nodalizację dla najważniejszej części instalacji, czyli grzałek symulujących pręty paliwowe. Komórki od 3 do 17 w rurze numer 20 stanowią reprezentację pęczka grzałek. Komórki pierwsza oraz druga reprezentują dolną komorę, a osiemnasta i dziewiętnasta górną komorę. Warunki przepływu wody chłodzącej ustalone są przez komponent FILL (numer 10), natomiast warunki ciśnienia w teście są kontrolowane przez komponent BREAK (numer 30). Grzane pręty reprezentowane są przez jedną strukturę cieplną. Aby zapewnić odpowiedni poziom mocy dla pęczka grzałek ze strukturą cieplną połączony jest komponent power, reprezentujący moc

całkowitą dostarczaną do struktury cieplnej. Komponent power pozwala określić między innymi osiowy oraz promieniowy rozkład mocy w strukturze cieplnej, jak również spadek mocy w czasie.

Opracowany model instancji testowej Flecht Seaset został poddany procesowi kwalifikacji. Podobnie jak dla instalacji Marviken ze względu na małe skomplikowanie modelu, również sam proces był uproszczony. W przypadku eksperymentu Flecht Seaset nie jest możliwe przeprowadzenie długotrwałych obliczeń stanu ustalonego, ze względu na same założenia testu. Test jest przeprowadzany jako symulacja fazy zalewania rdzenia, więc zakładane jest, że na początku testu pręty paliwowe są odsłonięte i nie są chłodzone. Nie jest więc możliwe utrzymanie stabilnych parametrów opisujących system, ponieważ bez zapewnienia chłodzenia (od czego zaczyna się eksperyment, a więc stan nieustalony) temperatura koszulki paliwowej zaczyna rosnąć, aż po kilkunastu sekundach obliczeń symulacja jest przerywana z informacją, że temperatura koszulki zbliża się do temperatury jej topienia. Jest to moment, w którym kończy się zakres stosowalności kodu, ponieważ zarówno właściwości materiałowe jak i równania cieplno-przepływowe są zwalidowane na zakres do temperatury uszkodzenia koszulki paliwowej. Kwalifikację przeprowadzono więc w warunkach stanu nieustalonego, czyli w warunkach panujących w testach przeprowadzonych w instalacji doświadczalnej. Do obliczeń wybrano dziesięć testów o różnych warunkach początkowych oraz założeniach dotyczących przepływu chłodziwa. Dla każdego z testów wykonano obliczenia referencyjne, tj. obliczenia z domyślnymi wartościami wszystkich parametrów modelowych, stosując zawsze tą samą nodalizację. Obliczenia różnią się, więc jedynie warunkami początkowymi (np. rozkład temperatury i ciśnienia) i brzegowymi. Wyniki kwalifikacji przedstawiono w tabeli 5.4.

Tabela 5.4. Kwalifikacja wyników obliczeń referencyjnych dla dziesięciu testów w instalacji eksperymentalnej Flecht Seaset

Test	Maksymalna temperatura koszulki paliwowej [K]			Czas schłodzenia najgorętszego pręta [s]			Czas całkowitego zalania rdzenia [s]		
	Dane eksp.	Obliczenia TRACE	Błąd	Dane eksp.	Obliczenia TRACE	Błąd	Dane eksp.	Obliczenia TRACE	Błąd
30619	985	1008	2.3%	293	244	-16.6%	572	549	-4.0%
30817	1085	1084	-0.1%	225	204	-9.3%	398	394	-1.0%
31108	1180	1231	4.3%	174	157	-9.4%	341	344	1.0%
31203	1303	1317	1.1%	255	241	-5.3%	425	450	5.9%
31302	1154	1195	3.6%	147	124	-15.5%	252	247	-1.7%
31504	1423	1443	1.4%	331	340	2.8%	593	690	16.4%
31701	1172	1157	-1.2%	63	72	15.0%	200	130	-34.7%
31805	1500	1496	-0.3%	399	384	-3.6%	690	751	8.9%
32013	1429	1395	-2.4%	255	283	11.2%	457	503	10.2%
32114	1426	1426	0.0%	425	411	-3.3%	634	700	10.4%

Dane eksperymentalne oraz wyniki obliczeń porównano dla trzech parametrów: maksymalnej temperatury koszulki paliwowej, czasu schłodzenia najgorętszego pręta w miejscu, gdzie osiowy rozkład mocy jest największy (czyli na wysokości 1.83 metra) oraz czasu całkowitego zalania rdzenia. Dwa pierwsze parametry opisują przebieg zmian temperatury koszulki paliwowej w czasie, natomiast ostatni opisuje przebieg propagacji frontu zalewania. Najlepszą zgodność otrzymano dla maksymalnej temperatury koszulki paliwowej, gdzie maksymalny błąd wynosił 4.3% dla testu 31108. Dalszy przebieg temperatury koszulki paliwowej w prawie połowie przypadków obliczeniowych odbiegał od danych eksperymentalnych – dla czterech testów czas schłodzenia najgorętszego pręta różnił się o ponad 10%, a dla dwóch kolejnych był blisko tej wartości. Podobnie dla czterech testów czas całkowitego zalania rdzenia odbiegał od danych eksperymentalnych, szczególnie dla testu 31701, gdzie błąd wynosi 35%. Dla pozostałych obliczeń błąd nie przekraczał 9%. Uzyskane wyniki są, więc w większości przypadków poprawne, wskazując, że zastosowana nodalizacja pozwala na uzyskanie wyników obliczeń zgodnych z danymi eksperymentalnymi. W związku z tym, że warunki początkowe i brzegowe w testach były różne, rozbieżności uzyskane dla niektórych testów będzie można zniwelować poprzez kalibrację parametrów modelowych.

Tabela 5.5. Warunki początkowe oraz brzegowe dla testów fazy zalewania w instalacji Flecht Seaset wraz z przypisaniem testów do zastosowania

Test	Prędkość przepływu [m/s]	Ciśnienie w górnej komorze [MPa]	Temperatura wody chłodzącej [K]	Grupa podobieństwa	Zastosowanie		
					Wybór parametrów / walidacja	Metoda Bayesowska	Metoda uczenia maszyn
30619	0.0389	0.134	309.15	II	Walidacja		
30817	0.0389	0.27	326.15	II	Wybór parametrów	MCMC	Zestaw testowy
31108	0.079	0.13	306.15	III	Wybór parametrów	Błąd modelowy	Zestaw uczący
31203	0.0384	0.28	325.15	II	Walidacja	Błąd modelowy	Zestaw uczący
31302	0.0765	0.28	325.15	III		MCMC	Zestaw testowy
31504	0.025	0.28	324.15	I		MCMC	Zestaw testowy
31701	0.155	0.28	326.15	III	Walidacja	Błąd modelowy	
31805	0.021	0.28	324.15	I		Błąd modelowy	
32013	0.0264	0.41	339.15	I	Wybór parametrów + Walidacja		
32114	0.025	0.28	398.15	I		Błąd modelowy	Zestaw uczący

W tabeli 5.5 przedstawiono warunki początkowe dla testów wykorzystywanych w obliczeniach wrażliwości służących do wyboru parametrów wejściowych oraz w obliczeniach kwantyfikacji niepewności wejściowych. Testy przypisano do grup podobieństwa na podstawie jakościowej oceny przebiegów temperatury koszulki paliwowej w eksperymencie oraz na podstawie prędkości przepływu wody chłodzącej w teście.

Metoda uczenia maszyn wymaga zastosowania innych testów do zbiorów uczącego i testowego, stąd wybrano po jednym teście z każdej grup do obu tych zbiorów. Następnie po jednym z testów z każdej grupy przypisano do walidacji wyników. W związku z tym testy wybrane do wyboru parametrów mogły się pokrywać częściowo z testami wybranymi do zastosowania w obliczeniach kwantyfikacji niepewności. Do obliczeń MCMC w ramach Bayesowskiej metody kwantyfikacji niepewności wybrano te same testy które przypisano do zbioru testowego. W ten sposób wyniki działania obu algorytmów będą mogły być bezpośrednio porównane na tych samych danych eksperymentalnych. Do trenowania metamodelu błędu modelowego wykorzystano wyniki obliczeń referencyjnych dla pięciu testów. Powyższy wybór był subiektywnie wykonany przez autora, aby przeprowadzić analizę dla różnych warunków które mogą panować w rdzeniu podczas fazy zalewania.

5.2.3. Wybór parametrów wejściowych (kroki 9 oraz 10)

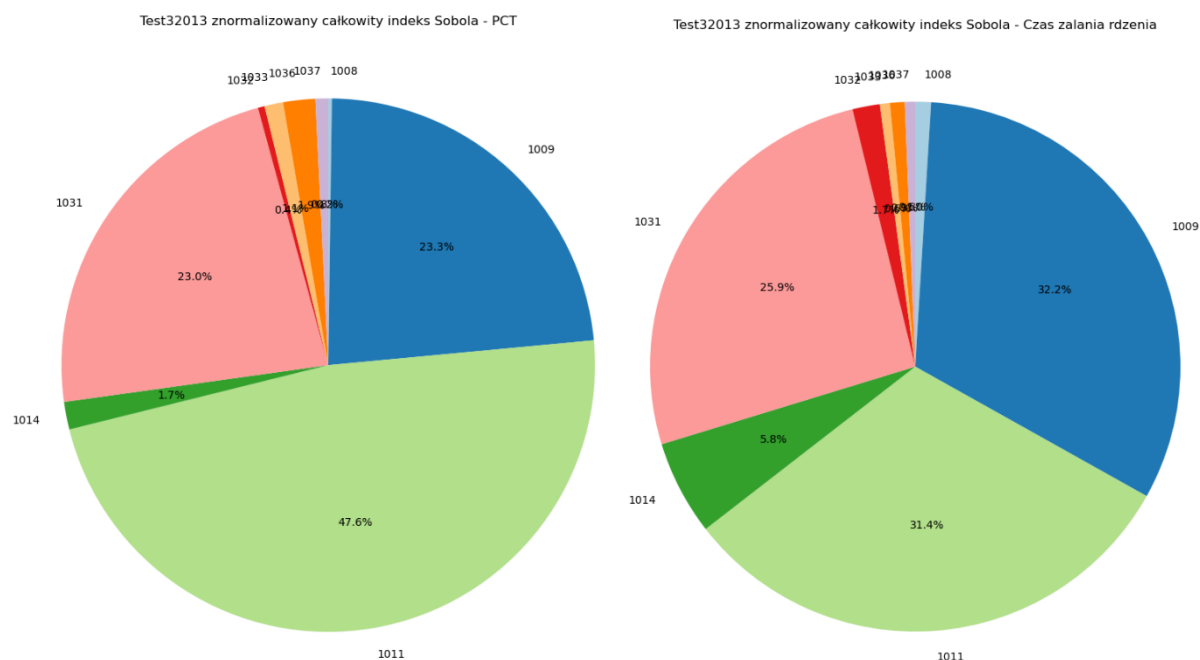
W kodzie TRACE dostępnych jest 38 parametrów wewnętrznych kodu, które zgodnie z opisem w dokumentacji kodu [50] można stosować do analiz wrażliwości i niepewności. Stanowią one zbiór mnożników oraz współczynników wykorzystywanych w obliczeniach cieplno-przepływowych. Domyślnie wartości dla wszystkich tych parametrów kalibracyjnych wynoszą jeden. Stąd w pierwszej kolejności wykonano lokalną analizę wrażliwości. Obliczenia wykonywano dla dwóch testów – 31701 oraz 31108. Obliczenia wykonywano poprzez zmianę w pliku wsadowym każdorazowo wartości tylko jednego parametru, tak aby bezpośrednio obserwować jego wpływ na wyniki przebiegu temperatury koszulki oraz poziomu wody w rdzeniu. Pierwszą serię obliczeń wykonano zmieniając wartość parametru z 1 na 100, a drugą serię obliczeń zmieniając wartość tego samego parametru z 1 na 0.1. Wyniki obliczeń porównano z obliczeniami referencyjnymi wykonanymi z wartościami domyślnymi parametrów. Na podstawie tak wykonanych obliczeń wybrano 15 parametrów, dla których zaobserwowano wpływ zmiany parametru na przebieg temperatury koszulki paliwowej oraz poziomu wody w rdzeniu. Następnie dokonano kolejnych obliczeń dla wartości parametrów równych 10, 2 oraz 0.5. Na podstawie analizy wyników tych obliczeń ograniczono po raz kolejny liczbę parametrów do dziewięciu, które wykazywały największy wpływ zmian ich

wartości na obserwowane wartości wyjściowe. Po wstępnej lokalnej analizie wrażliwości, dla wybranych dziewięciu parametrów wykonano obliczenia globalnej analizy wrażliwości stosując opracowaną metodą opisaną w [rozdziale 4.4](#).

Metoda globalnej analizy wrażliwości pozwala na ocenę wpływu zmian wartości parametrów na wariancję wielkości wyjściowych. Następnie można dokonać rankingu istotności parametrów, na podstawie analizy całkowitych indeksów Sobola które są miernikiem wpływu parametru na wariancję wielkości wyjściowej. Na potrzeby globalnej analizy wrażliwości określono sześć skalarnych wielkości wyjściowych. Trzy opisujące przebieg temperatury koszulki paliwowej oraz trzy opisujące zmiany poziomu wody w rdzeniu. Temperaturę koszulki paliwowej badano na wysokości 1.98 metra, gdzie gęstość mocy była największa. Dla danych eksperymentalnych wykorzystano pomiary z termopary na tej wysokości, natomiast dla obliczeń była to wartość z komórki struktury cieplnej której środek znajdował się również na wysokości 1.98 metra. Dla temperatury koszulki paliwowej określono następujące skalarne wielkości wyjściowe – maksymalna temperatura koszulki paliwowej (dalej oznaczenie - PCT), czas, w którym osiągnięto maksymalną temperaturę koszulki paliwowej (dalej oznaczenie - czas PCT) oraz czas, w którym osiągnięto minimalną wartość temperatury, czyli czas schłodzenia najgorętszego pręta paliwowego (dalej oznaczenie - czas temp min). Dla przebiegu poziomu wody w rdzeniu określono następujące wielkości wyjściowe – czas, gdy rdzeń jest całkowicie zalany, czyli woda sięga na wysokość 3.66 metra (dalej oznaczenie czas zalania), wartość poziomu wody w rdzeniu w 1/3 czasu trwania eksperymentu oraz wartość poziomu wody w rdzeniu 2/3 czasu trwania eksperymentu.

Stosując generator wartości losowych Saltelliego wykonano 2560 obliczeń dla każdego z trzech testów oznaczonych w tabeli 5.5 etykietą „Wybór parametrów”, czyli testu 30817, 31108 oraz 32013. Obliczenia wykonano zmieniając wartości dziewięciu parametrów wewnętrznych kodu, wartości losowano każdorazowo z rozkładu jednorodnego o zakresie od 0.1 do 5. Po wykonaniu obliczeń z zastosowaniem kodu TRACE dokonano analizy rezultatów stosując metody zaimplementowane w bibliotece SALiB do wyznaczenia indeksów Sobola pierwszego rzędu oraz całkowitych indeksów Sobola. Wyniki obliczeń znormalizowanych całkowitych indeksów Sobola dla testu 32013 (wybranego jako przykładowy) oraz wielkości wyjściowych – maksymalnej temperatury koszulki (po lewej) i czasu zalania rdzenia (po prawej) przedstawiono na rysunku 5.20. Każdy wycinek koła prezentuje procentowo wpływ danego parametru wejściowego na wariancję wielkości wyjściowej uwzględniając efekty pierwszego rzędu oraz wyższych rzędów. Zmiany parametrów o oznaczenia 1009, 1011 oraz 1031

posiadają w tym przypadku największy wpływ na uzyskiwane wartości wielkości wyjściowych. Wartości znormalizowanych całkowitych indeksów Sobola dla sześciu analizowanych wielkości wyjściowych przedstawiono w tabeli 5.6. Wartość indeksu Sobola dla każdego z parametrów dla danej wielkości wyjściowej powstała z uśrednienia trzech całkowitych indeksów Sobola obliczonych dla każdego z testów. Ostatni wiersz w tabeli przedstawia uśrednienie wartości całkowitych indeksów Sobola dla wszystkich sześciu badanych wielkości wyjściowych.



Rys. 5.20. Wyniki obliczeń całkowitego indeksu Sobola stosując globalną analizę wrażliwości dla instalacji eksperymentalnej Flecht Seaset dla dwóch wybranych wielkości wyjściowych

Tabela 5.6. Wyniki globalnej analizy wrażliwości dla zjawiska zalewania rdzenia – znormalizowane całkowite indeksy Sobola dla dziewięciu badanych parametrów wejściowych kodu

Wielkość wyjściowa	Całkowity indeks Sobola	Oznaczenie parametru wejściowego								
		1008	1009	1011	1014	1031	1032	1033	1036	1037
PCT [K]	S_{T1}	0.002	0.180	0.601	0.015	0.176	0.004	0.006	0.013	0.004
Czas PCT [s]	S_{T2}	0.021	0.170	0.489	0.052	0.142	0.021	0.037	0.048	0.020
Czas temp min [s]	S_{T3}	0.229	0.115	0.180	0.137	0.107	0.013	0.046	0.130	0.044
Czas zalania [s]	S_{T4}	0.012	0.175	0.354	0.116	0.273	0.017	0.013	0.028	0.013
Poziom wody w 1/3 czasu [m]	S_{T5}	0.019	0.083	0.171	0.254	0.177	0.033	0.079	0.116	0.070

Wielkość wyjściowa	Całkowity indeks Sobola	Oznaczenie parametru wejściowego								
		1008	1009	1011	1014	1031	1032	1033	1036	1037
Poziom wody w 2/3 czasu [m]	S_{T6}	0.025	0.126	0.264	0.105	0.323	0.041	0.029	0.058	0.030
	$\bar{S}_T$	0.051	0.141	0.343	0.113	0.200	0.022	<u>0.035</u>	0.065	0.030

Całkowite indeksy Sobola są znormalizowane więc sumują się do wartości równej jeden, stąd dzieląc tą wartość przez liczbę parametrów (dziewięć) otrzymuje się średnią wartość indeksu Sobola, gdyby parametry miały taki sam wpływ, równą 0.11. Każdorazowo, gdy wartość indeksu Sobola przekraczała tą wartość zaznaczano ten fakt w tabeli. Jeden parametr – współczynnik wymiany ciepła dla w rozproszonego przepływu kropelkowego (ang. „dispersed flow film boiling heat transfer coefficient”) oznaczony numerem **1011** miał każdorazowo wartość indeksu Sobola przekraczającą wartość średnią. Uśredniony całkowity indeks Sobola dla tego parametru również był najwyższy i wynosi 34.3%. Dwa kolejne parametry – współczynnik wymiany ciepła między parą a ścianką (ang. „single phase vapor to wall heat transfer coefficient”) oznaczony numerem **1009** oraz opór hydrodynamiczny między fazami w warunkach rozproszonego przepływu kropelkowego (ang. „interfacial drag coefficient for dispersed flow film boiling”) oznaczony numerem **1031** – w pięciu przypadkach charakteryzowały się wartością indeksu Sobola większą od średniej. Ostatecznie były to parametry o drugiej i trzeciej największej wartości uśrednionego całkowitego indeksu Sobola. Ostatnim parametrem wybranym do dalszej analizy na podstawie globalnej analizy wrażliwości był mnożnik dla krytycznego strumienia przepływu oznaczony numerem **1014**. Dodatkowo na podstawie wykonanej wcześniej lokalnej analizy wrażliwości dodano do dalszej analizy parametr określający opór hydrodynamiczny między fazami w warunkach odwróconego przepływu pierścieniowego (ang. „interfacial drag coefficient for inverted annular flow”) oznaczony numerem **1033**. Wpływ tego parametru był wyraźnie obserwowalny przy wykonywaniu lokalnej analizy wrażliwości. Dodatkowo biorąc pod uwagę występowanie tego typu przepływu w trakcie fazy zalewania rdzenia w czasie awarii typu LOCA uznano, że parametr ten warto uwzględnić w dalszej analizie. Ostatecznie do dalszej analizy wybrano, więc pięć parametrów modelowych dostępnych w kodzie TRACE. Ze względu na prostotę zapisu w dalszej części pracy będą używane oznaczenia numeryczne dla wybranych parametrów.

5.2.4. Kwantyfikacja niepewności parametrów wejściowych (krok 11)

Podobnie jak dla określania niepewności parametrów wejściowych modelu wpływu krytycznego na podstawie obliczeń dla instalacji Marviken, dla parametrów charakteryzujących fazę zalewania rdzenia w trakcie awarii LBOLCA dokonano kwantyfikacji niepewności

wybranych parametrów stosując dwie opracowane metody. Zgodnie z przyporządkowaniem testów wykonanych w instalacji Flecht Seaset obliczenia dla zestawu uczącego w metodzie uczenia maszyn oraz obliczenia stosując algorytm MCMC wykonano dla tego samego zbioru testów. Testy te o oznaczeniach 30817, 31302, 31504 posiadały różne warunki, co przekładało się na rozbieżności w przebiegach poziomu wody oraz temperatury koszulki. Dodatkowo, przykładowo dla testu 31504 przebieg obliczeń referencyjnych dla temperatury koszulki paliwowej był zbliżony z danymi eksperymentalnymi, natomiast poziom wody w rdzeniu nie odzwierciedlał poprawnie danych pomiarowych. Tym samym zmiany parametrów mogą powodować poprawę jednej wielkości wyjściowej, jednocześnie pogarszając wyniki dla drugiej. Dla każdego z testów algorytmy zaimplementowane w obu metodach powinny więc dążyć do zapewnienia odpowiedniej równowagi, tak aby przebiegu obu wielkości wyjściowego dla jednego zestawu parametrów były na akceptowalnym poziomie. Różnorodność warunków początkowych oraz przebiegów badanych wielkości wyjściowych w wybranych testach z jednej strony zapewnia uniwersalność uzyskanych ostatecznie rozkładów gęstości prawdopodobieństwa, ale z drugiej strony istnieje ryzyko, że wartości parametrów dające najlepsze wyniki w jednym teście mogą nie być odpowiednie dla innego testu. W czasie awarii typu LBLOCA mogą jednak panować różne warunki przepływu wody z systemów bezpieczeństwa, stąd podobna niepewność jest również wbudowana w wykonywanie obliczeń dla reaktorów energetycznych, gdzie wyników nie można porównać z danymi eksperymentalnymi.

5.2.4.1. Kwantyfikacja niepewności parametrów wyjściowych stosując metodę opartą o podejście Bayesowskie

Przed rozpoczęciem obliczeń MCMC w ramach metody opartej o podejście Bayesowskie należy określić metamodel dla rozbieżności modelowej, pierwsze oszacowanie wartości parametrów modelowych, rozkład propozycji, rozkład a’priori parametrów oraz liczbę obliczeń i liczbę odrzucanych obliczeń.

Ze względu na badanie dwóch wielkości wyjściowych – temperatury koszulki paliwowej oraz poziomu wody w rdzeniu – konieczne było przygotowanie metamodeli dla rozbieżności modelowej dla obu tych wielkości. Dokonano więc porównania wyników obliczeń referencyjnych z danymi eksperymentalnymi dla pięciu testów oznaczonych etykietą błąd modelowy w tabeli 5.5. Zestaw testów wybranych do opracowania metamodelu obejmował testy o najdłuższym czasie trwania (800 sekund – test 31805) oraz najkrótszym czasie trwania (200 sekund – test 31701), test o największej maksymalnej temperaturze koszulki paliwowej

(1500 K – test 31805), jak również testy o największych i najmniejszych wartościach błędów dla trzech wielkości przedstawionych w tabeli 5.4. Testy były więc zróżnicowane, tak aby dane wejściowe dla metamodelu obejmowały zarówno testy o dobrej zgodności obliczeń referencyjnych z danymi pomiarowymi jak i testy o słabej zgodności. Dane wyjściowe z metamodelu stanowił zestaw 2000 procesów Gaussowskich reprezentujących przebieg rozbieżności modelowej dla temperatury koszulki paliwowej oraz poziomu wody w rdzeniu.

Pierwszym oszacowaniem dla parametrów modelowych o oznaczeniach 1009, 1011, 1014, 1031 oraz 1031 były ich wartości domyśle wynoszące 1. W pierwszym kroku algorytmu MCMC nowa propozycja parametrów jest losowana z rozkładu propozycji a współczynnik Metropolisa (równanie 4.10) obliczany w oparciu o porównaniu funkcji wiarygodności dla obliczeń wykonanych z nową propozycją wartości parametrów a obliczeniami referencyjnymi, wykonanymi z wartościami domyślnymi parametrów (czyli z pierwszym oszacowaniem). Jako rozkład propozycji wybrano wielowymiarowy rozkład normalny (liczba wymiarów równa jest liczbie parametrów, w tym wypadku to pięć wymiarów), gdzie każdy z rozkładów posiadał stałe odchylenie standardowe równe 0.2 oraz wartość średnią przyjmującą wartość parametru w poprzednim kroku. W pierwszym kroku losowano więc z rozkładów o wartości średniej równej 1.0 i odchyleniu standardowym równym 0.2.

Rozkład a’priori dla każdego z parametrów był określony jednakowo jako rozkład jednorodny w zakresie od 0 do 10. Wartość 10 została wybrana na podstawie wykonanej lokalnej analizy wrażliwości. Maksymalna temperatura koszulki paliwowej przy takiej wartości dla parametrów 1014 oraz 1031 wzrastała o ponad 20%, więc w przypadku testu 31504 taki wzrost maksymalnej temperatury koszulki powodowałby przerywanie obliczeń poprzez przekroczenie zakresu stosowności kodu. Prawdopodobieństwo a’priori dla propozycji, w której jedna z wartości przekroczyłaby 10 wynosiło więc zero. Tym samym rozkład a’posteriori również będzie ograniczony w zakresie między 0 a 10.

Liczba obliczeń wykonywana dla instalacji eksperymentalnej Flecht Seaset była zdecydowanie mniejsza niż dla instalacji Marviken. Wynika to z faktu, że obliczenia dla instalacji Flecht Seaset były znacznie bardziej czasochłonne niż dla instalacji Marviken. Obliczenia dla instalacji w których występuje wymiana ciepła, oraz wzrasta temperatura koszulki paliwowej zajmują więcej czasu niż obliczenia w których kod cieplno-przepływowy nie uwzględnia zjawisk związanych z wymianą ciepła (tak jak dla instalacji Marviken). Jest to szczególnie widoczne w obliczeniach, gdy temperatura koszulki paliwowej wzrasta powyżej 1500 K. Wówczas pojedyncze obliczenia mogą trwać nawet godzinę, w porównaniu do standardowych

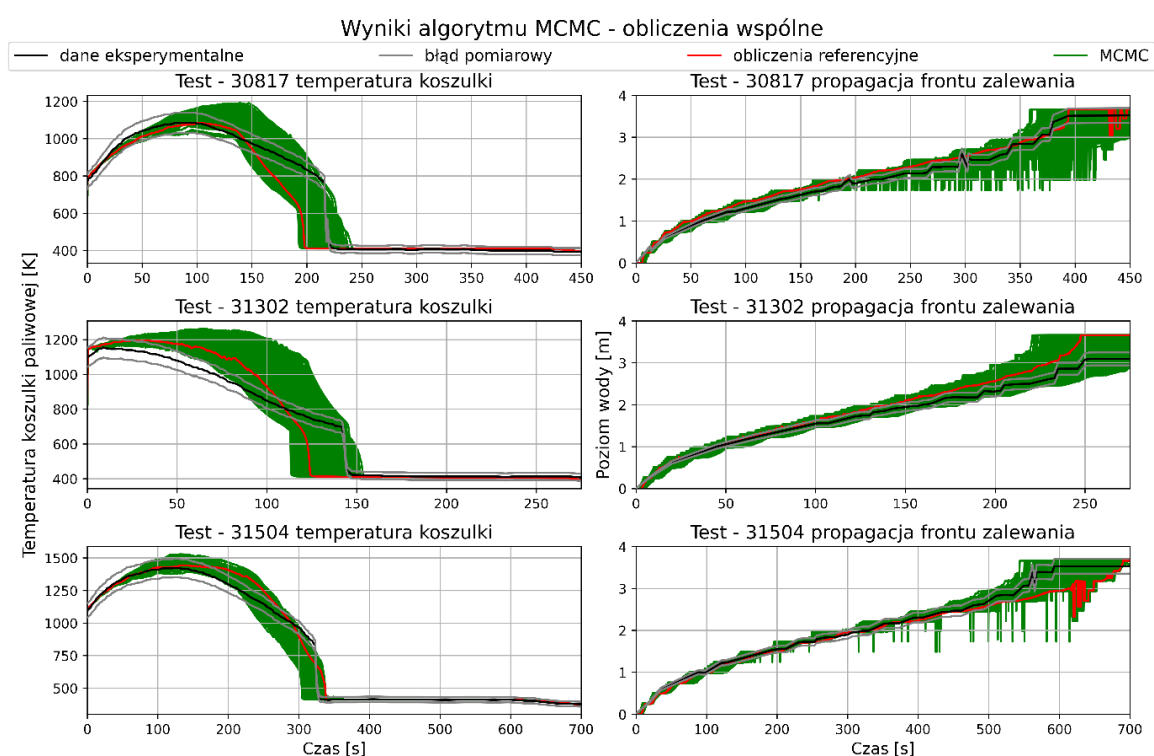
czasów wykonywania obliczeń wynoszących około dwóch minut. Dodatkowo czasy trwania eksperymentów były zdecydowanie dłuższe – średnio około 500 sekund w porównaniu do 70 sekund dla instalacji Marviken, co również przekładało się na dłuższy czas wykonywania obliczeń. Dla porównania obliczenia referencyjne dla jednego testu w instalacji Marviken trwały od dwóch do pięciu sekund, natomiast dla instalacji Flecht Seaset od 20 do nawet 180 sekund. Tym samym wykonywanie obliczeń było znacznie bardziej czasochłonne.

Obliczenia wykonano stosując dwa podejścia. W pierwszym podejściu wykonywano obliczenia stosując algorytm MCMC i obliczając wpierw funkcję wiarygodności dla temperatury koszulki paliwowej, a następnie dla poziomu wody w rdzeniu dla każdej z propozycji. Współczynnik Metropolisa obliczano więc dwa razy – oddzielnie dla każdej z badanych wielkości wyjściowych. Propozycja parametrów była akceptowana tylko gdy propozycję zaakceptowano dla obu wielkości wyjściowych. Dla każdego z trzech testów wykonano 12 000 obliczeń. W drugim podejściu wykonywano oddzielne obliczenia w których algorytm obliczał funkcję wiarygodności tylko dla jednej z badanych wielkości. Wykonano w ten sposób po 8 000 obliczeń dla temperatury koszulki paliwowej oraz poziomu wody w rdzeniu. Dla jednego testu wykonywano, więc łącznie 16 000 obliczeń. W obu podejściach pierwsze 10% obliczeń odrzucane było z dalszej analizy. Obliczenia wykonywano stosując dwa podejścia, ponieważ w literaturze nie odnaleziono zaleceń w jaki sposób przeprowadzać obliczenia stosując algorytm MCMC dla dwóch badanych wielkości wyjściowych. Równocześnie uznano, że takie podejście umożliwi porównanie uzyskanych rozkładów gęstości prawdopodobieństwa dla parametrów modelowych dla wielkości wyjściowych badanych oddzielnie oraz wspólnie. Pozwoli to również na sformułowanie rekomendacji dla dalszych badań wykonywanych stosując algorytm MCMC dla więcej niż jednej badanej wielkości wyjściowej. W dalszej części rozprawy pierwsze podejście będzie określane jako obliczenia wspólne natomiast drugie jako obliczenia oddzielne.

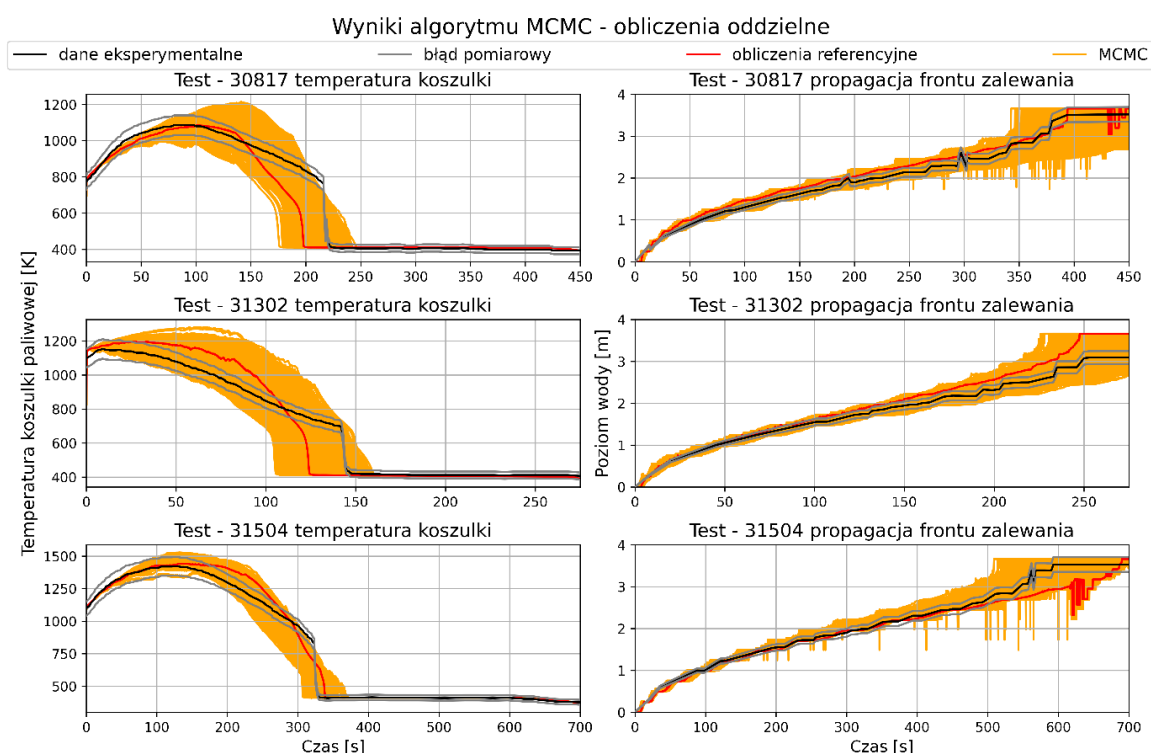
Tabela 5.7. Liczba zaakceptowanych propozycji w algorytmie MCMC dla obliczeń dla instalacji Flecht Seaset, wraz z procentowym udziałem zaakceptowanych propozycji w odniesieniu do całkowitej liczby obliczeń w podziale na dwa podejścia

Liczba zaakceptowanych propozycji						
	Obliczenia wspólne		Obliczenia oddzielne			
	Obie wielkości	Udział procentowy	Temperatura koszulki	Udział procentowy	Poziom wody	Udział procentowy
Test 30817	5 499	46 %	4 233	53 %	5 564	70 %
Test 31302	7 394	62 %	5 668	71 %	5 887	74 %
Test 31504	3 261	27 %	3 432	43 %	5 262	66 %
Suma	16 154	45 %	13 333	56 %	16 713	70 %

W tabeli 5.7 zaprezentowano liczbę zaakceptowanych przez algorytm MCMC propozycji wartości parametrów w przypadku obliczeń wspólnych oraz oddzielnych dla trzech testów, w których badano zjawisko zalewania rdzenia w instalacji eksperymentalnej Flecht Seaset. Dla każdego testu procentowy udział zaakceptowanych propozycji był mniejszy dla obliczeń wspólnych. Wskazuje to, że uzyskanie wartości parametrów wejściowych które pozwolą na uzyskanie dobrej zgodności obliczeń z eksperymentem dla obu badanych wielkości wyjściowych jest trudniejsze niż w przypadku badania tylko jednej wielkości. W przypadku obliczeń oddzielnych każdorazowo więcej obliczeń było akceptowanych dla obliczeń w których badaną wyjściową wielkością był poziom wody w rdzeniu (propagacja frontu zalewania). Wskazuje to, że przebieg temperatury koszulki jest bardziej wrażliwy na zmiany parametrów wejściowych – zmiany w parametrach mają większy wpływ na różnicę w funkcjach wiarygodności pomiędzy propozycją parametrów, a zestawem wartości parametrów poprzednio zaakceptowanym.



Rys. 5.21. Wyniki obliczeń wspólnych dla zaakceptowanych przez algorytm MCMC propozycji wartości parametrów dla każdego z testów, w porównaniu z danymi eksperymentalnymi i obliczeniami referencyjnymi dla obu badanych wielkości wyjściowych



Rys. 5.22. Wyniki obliczeń oddzielnych dla dwu badanych wielkości wyjściowych, dla zaakceptowanych przez algorytm MCMC propozycji wartości parametrów dla każdego z testów, w porównaniu z danymi eksperymentalnymi i obliczeniami referencyjnymi

Na rysunkach 5.21 oraz 5.22 zaprezentowano przebiegi temperatury koszulki paliwowej oraz poziomu wody w rdzeniu dla wszystkich zaakceptowanych propozycji parametrów modelowych dla obliczeń wspólnych (rysunek 5.21) oraz obliczeń oddzielnych (rysunek 5.22). Wyniki obliczeń zestawiono z danymi eksperymentalnymi oraz obliczeniami referencyjnymi (z domyślnymi wartościami parametrów modelowych). W przypadku przebiegów temperatury koszulki zarówno dla obliczeń wspólnych jak i oddzielnych zakres uzyskanych wyników pokrywał dane eksperymentalne dla testów 30817 oraz 31504. W przypadku testu 31302 tylko obliczenia oddzielne pozwoliły na pokrycie danych eksperymentalnych w całym zakresie czasowym. Jednocześnie rozrzut wyników jest większy w przypadku obliczeń oddzielnych, uzyskane zakresy są więc szersze. W przypadku testów 30817 oraz 31302 można zauważyć przesunięcie maksymalnej temperatury koszulki w czasie w porównaniu do danych eksperymentalnych. Wynika to z zasady obliczenia funkcji wiarygodności oraz charakterystyki przebiegu temperatury koszulki paliwowej. Przebieg ten charakteryzuje się istnieniem momentu w czasie, gdy temperatura koszulki osiąga w sposób gwałtowny temperaturę nasycenia dla ciśnienia w obudowie bezpieczeństwa, czyli rdzeń na tej wysokości zostaje zalany. Temperatura spada wówczas w ciągu kilku sekund o 300 – 400 °C (w zależności od

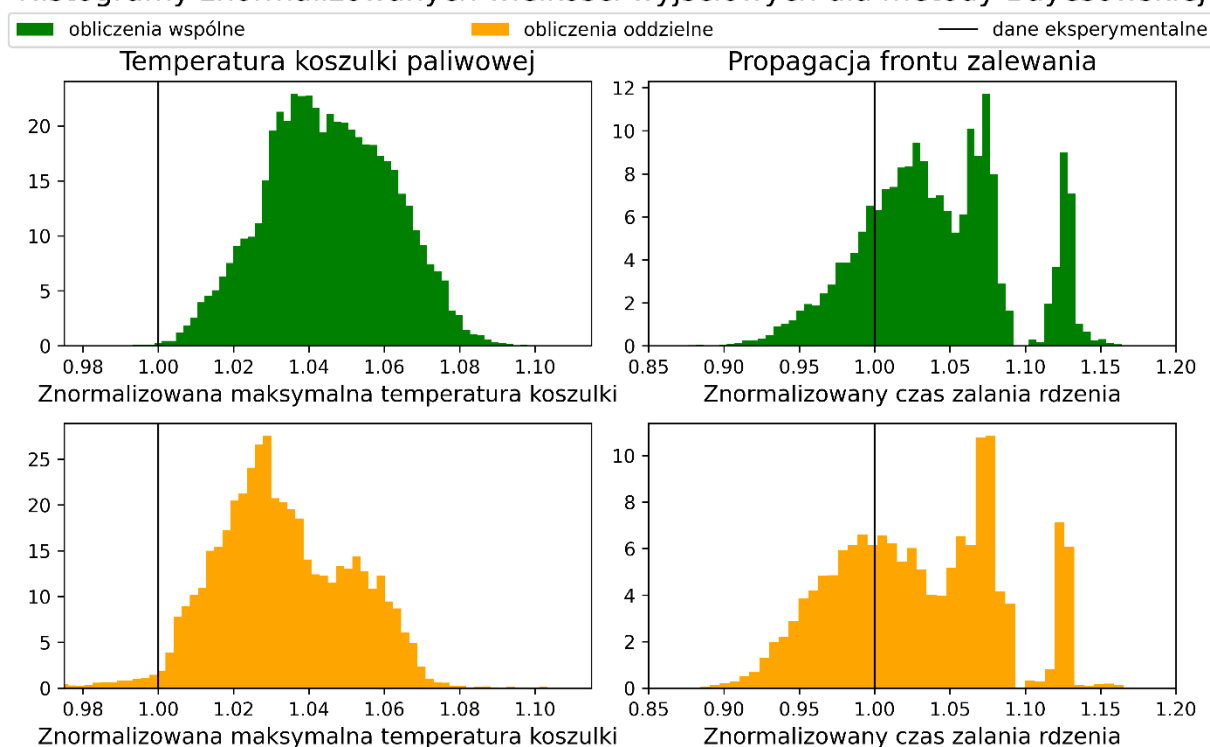
testu). Stąd, gdy w obliczeniach ten moment czasowy zostanie wyznaczony zbyt wcześnie, wówczas cały zakres czasowy aż do momentu spadku temperatury w eksperymencie będzie się charakteryzował bardzo dużymi różnicami między danymi eksperymentalnymi a wartościami obliczeniowymi. Funkcja wiarygodności będzie większa gdy maksymalna temperatura koszulki będzie przez kilkanaście-kilkadziesiąt sekund większa od eksperymentalnej o 100 °C niż w przypadku gdy do zalania rdzenia na badanej wysokości dojdzie o kilkanaście-kilkadziesiąt sekund za wcześnie i różnice między danymi eksperymentalnymi, a wynikami obliczeń będą wynosić od 300 do nawet 500 °C. Algorytm będzie więc akceptować propozycje w których moment czasowy zalania rdzenia dla obliczeń będzie zbliżony do danych eksperymentalnych. Stąd wynika przesuwanie wartości maksymalnej temperatury koszulki oraz jej większe wartości. Dodatkowo na podstawie danych w tabeli 5.5 oraz ilustracji przebiegu temperatury koszulki paliwowej można zauważyć prawidłowość, że im obliczenia referencyjne były bardziej zbieżne z danymi eksperymentalnymi tym mniej propozycji było akceptowanych. W przypadku testu 31302 gdzie obliczenia referencyjne najbardziej odbiegały od danych eksperymentalnych, algorytm eksplorując przestrzeń możliwych wartości parametrów uzyskiwał coraz lepsze dopasowanie do danych eksperymentalnych. Natomiast w przypadku testu 31504 obliczenia referencyjne mało odbiegały od eksperymentu, stąd algorytm akceptował małą liczbę propozycji, które faktycznie były zbieżne z danymi eksperymentalnymi. Stąd można wnioskować, że dla testu 31302 procentowy udział odrzucanych wartości mógłby być większy niż założone 10%.

W przypadku propagacji frontu zalewania nie zaobserwowano znacznych różnic między obliczeniami wspólnymi a obliczeniami oddzielnymi. We wszystkich przypadkach obliczenia pokrywają dane eksperymentalne a uzyskane zakresy niepewności nie są nadmiernie szerokie. W obliczeniach oddzielnych różnice między liczbą zaakceptowanych propozycji dla każdego z testów są rzędu kilku procent, więc nie zaobserwowano zależności między liczbą zaakceptowanych propozycji, a jakością obliczeń referencyjnych, którą opisano dla przebiegu temperatury koszulki paliwowej.

W celu przeprowadzenia ilościowej analizy otrzymanych rezultatów działania algorytmu porównano wartości maksymalnej temperatury koszulki oraz czasu zalania rdzenia dla każdego z zaakceptowanych obliczeń w odniesieniu do odpowiadających wielkości eksperymentalnych dla każdego testu. Wyniki dla wszystkich testów znormalizowano w odniesieniu do danych eksperymentalnych i stworzono histogramy znormalizowanych wartości dla obliczeń wspólnych oraz obliczeń oddzielnych, co przedstawiono na rysunku 5.23. Rozbieżności

maksymalnej temperatury koszulki paliwowej zaakceptowanych propozycji w zdecydowanej większości nie przekroczyły 8% w odniesieniu do danych eksperymentalnych. W przypadku czasu zalania rdzenia rozbieżności były większe, sięgały maksymalnie 15%, jednak większość wyników mieściła się w granicach rozbieżności wynoszącej 10%. Nie zaobserwowano również znaczących różnic w wynikach w zależności od wykonywania obliczeń funkcji wiarygodności. Algorytm działa więc poprawnie, nie pozwalając na zaakceptowanie propozycji parametrów które skutkowałyby wykonaniem obliczeń charakteryzującymi się znaczącymi rozbieżnościami z danymi eksperymentalnymi.

Histogramy znormalizowanych wielkości wyjściowych dla metody Bayesowskiej



Rys. 5.23. Znormalizowana maksymalna temperatura koszulki paliwowej oraz maksymalny poziom wody dla wszystkich zaakceptowanych propozycji algorytmu MCMC dla testów 30817, 31302 oraz 31504 w instalacji Flecht Seaset

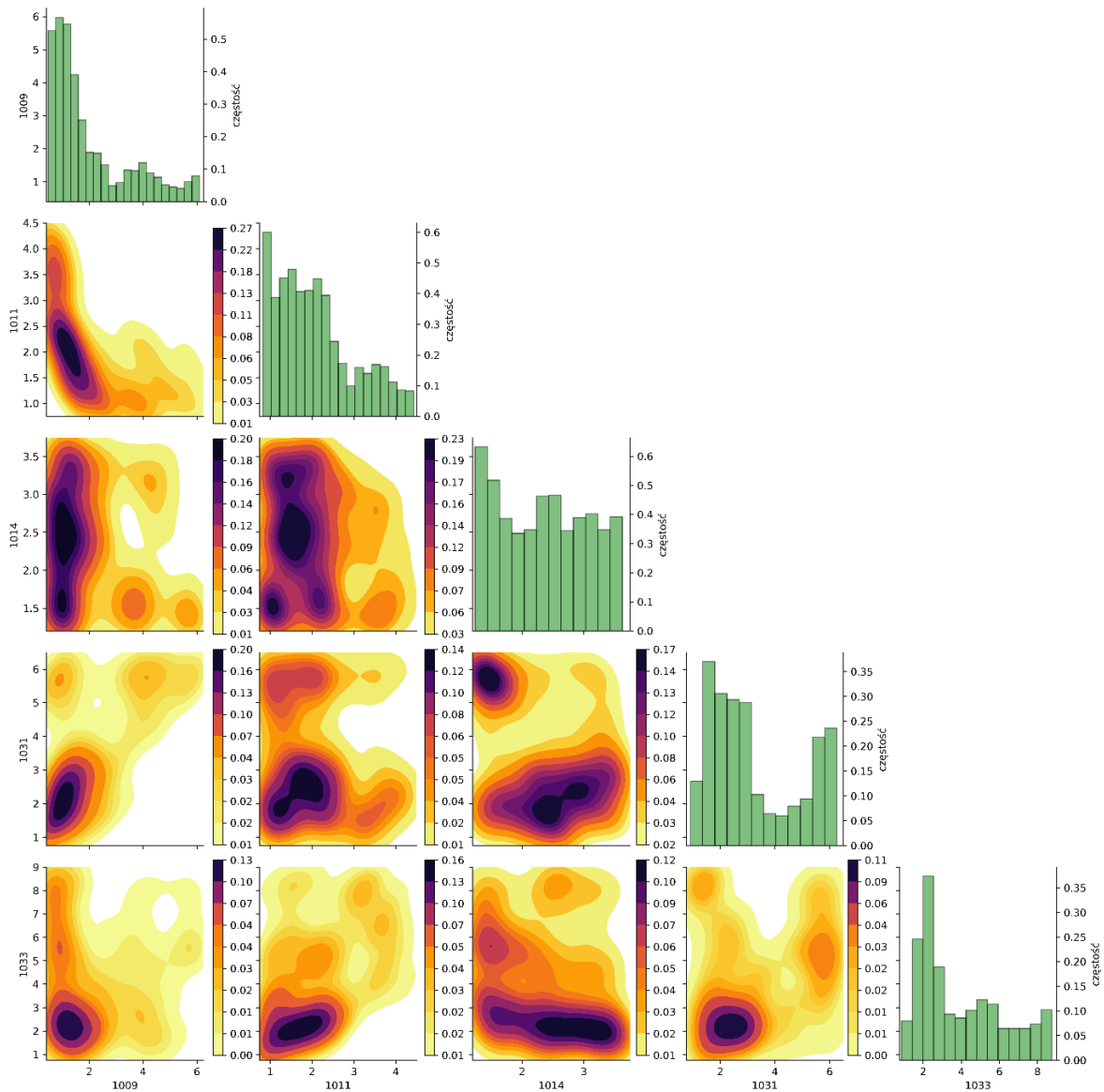
Określenie rozkładów gęstości prawdopodobieństwa parametrów

Podobnie jak w przypadku określenia rozkładów gęstości prawdopodobieństwa parametrów modelowych kodu dla instalacji Marviken, tak i dla wyników z instalacji eksperymentalnej Flecht Seaset w pierwszej kolejności dokonano obróbki zbiorów parametrów uzyskanych dla każdego testu. Dla obliczeń wspólnych otrzymano, więc trzy rozkłady dla każdego z pięciu parametrów (po jednym dla danego testu). Dla większości parametrów rozkłady te pokrywały się ze sobą jedynie fragmentarycznie, więc dokonano w pierwszej kolejności ich połączenia poprzez zastosowanie

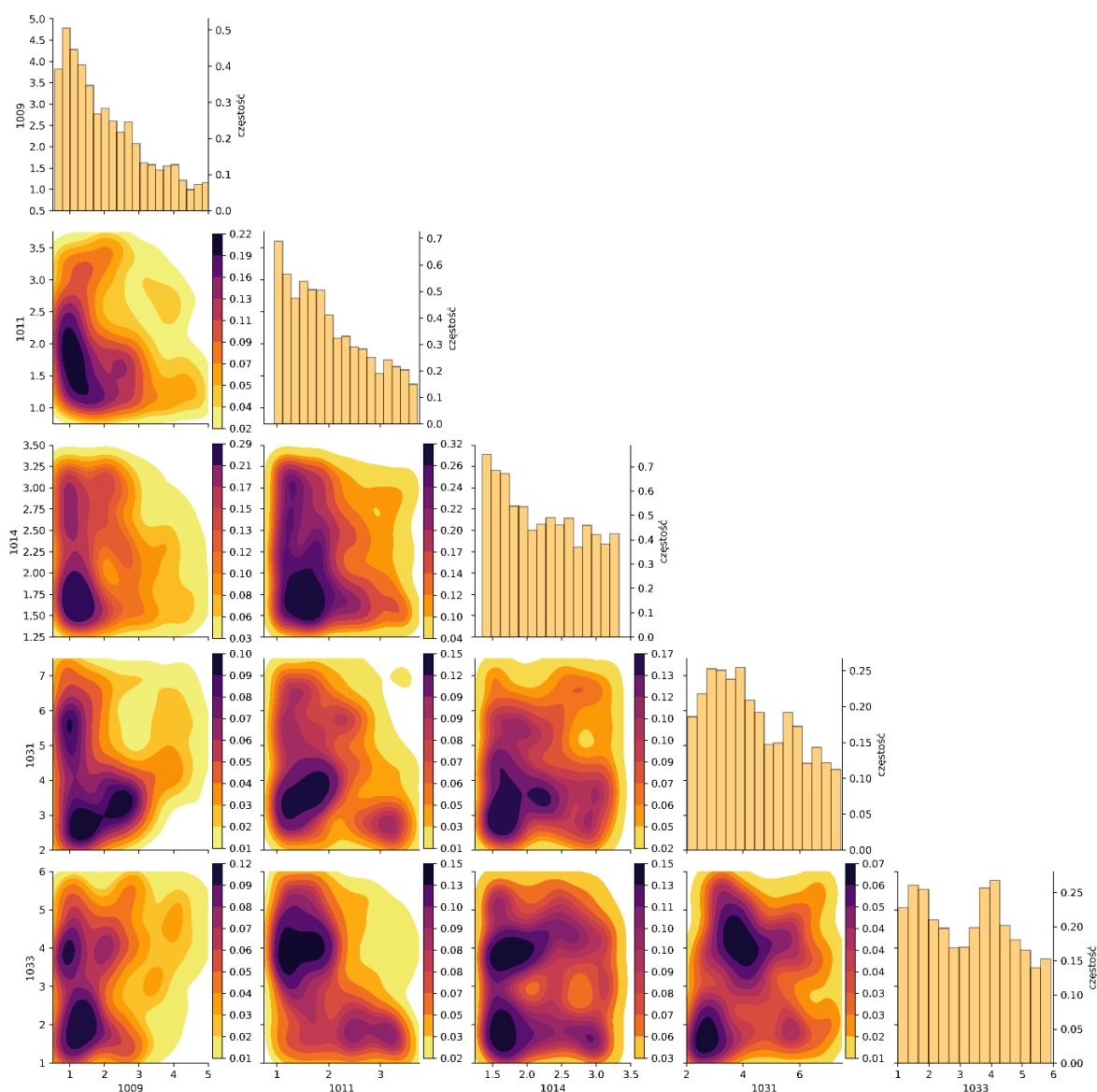
iloczynu logicznego, a następnie uśrednienia. Aby wykonać operację iloczynu logicznego na rozkładach określono wartość odpowiadającą 2.5 percentyla oraz 97.5 percentyla. Odrzucono więc górne 2.5% i dolne 2.5%. Następnie dla tak wyznaczonych granic dla każdego parametru wyznaczono wartość maksymalną z 2.5 percentyla oraz minimalną z 97.5 percentyla z trzech testów. W ten sposób w tak określonym zakresie zmienności parametrów znajdują się tylko wartości parametrów które pozwalały na uzyskanie akceptowalnych propozycji w każdym z testów. Ostateczne rozkłady dla każdego z parametrów są, więc kompromisem między rozkładami uzyskanymi oddzielnie dla każdego testu i nie zawsze odzwierciedlają lokalnych maksimów jakie zidentyfikowano dla pojedynczych rozkładów. Różnice w indywidualnych rozkładach wynikają z omówionych wcześniej różnic w akceptacji propozycji oraz zgodności obliczeń referencyjnych z danymi eksperymentalnymi. Dla obliczeń oddzielnych początkowo należało połączyć rozkłady uzyskane dla temperatury koszulki paliwowej oraz rozkłady dla propagacji frontu zalewania dla wszystkich parametrów w poszczególnych testach. Ponownie poszukiwane były części wspólnych tych rozkładów, aby znaleźć wartości parametrów które pozwolą na uzyskanie akceptowalnych wyników jednocześnie dla obu analizowanych wielkości wyjściowych. Uzyskano więc w ten sposób rozkłady dla pięciu parametrów dla trzech testów. Dalsza procedura przebiegała, więc identycznie jak dla obliczeń wspólnych. Stosując podejście oparte o obliczenia oddzielne należy, więc dodać jeden dodatkowy krok obróbki wyników. Porównanie rozkładów uzyskiwanych dla poziomu wody w rdzeniu oraz temperatury koszulki paliwowej wskazywało, że mimo że rozkłady te często się od siebie różniły, to nie odnotowano problemów z wydzielaniem znaczącej części wspólnej. Uzyskane rozkłady parametrów zaprezentowano na rysunku 5.24 dla obliczeń wspólnych oraz na rysunku 5.25 dla obliczeń oddzielnych. Ponownie rozkłady zaprezentowano w formie histogramów dla pojedynczych parametrów oraz dwuwymiarowych rozkładów dwóch zmiennych dla wszystkich możliwych kombinacji. Dla histogramów oś y przedstawiona jest po ich prawej stronie. Oś y umiejscowiona po lewej stronie stosowana jest dla dwuwymiarowych rozkładów par parametrów.

Dla parametru 1009 w obu przypadkach najbardziej prawdopodobne są wartości między 0.5 a 1.75, następnie częstość wartości parametru maleje wykładniczo. W obliczeniach wspólnych górna granica zakresu prawdopodobnych wartości wynosi 6.1 natomiast w przypadku obliczeń oddzielnych 5. Dla parametru 1011 ponownie największe częstości występują w zakresie zawierającym wartość domyślną między 0.8 a 2.0. Dalej podobnie jak dla parametru 1009 częstości akceptowanych wartości parametru maleją wykładniczo do 4.45 dla obliczeń

wspólnych oraz 3.7 dla obliczeń oddzielnych. Rozkłady dla parametru 1014 dla obu metod obliczeniowych rozkłady są bardzo podobne. Ograniczone są w zakresach od 1.24 do 3.63 w obliczeniach wspólnych oraz od 1.35 do 3.3 w obliczeniach oddzielnych. Wartości około 1.5 są najbardziej prawdopodobne, pozostałe wartości w określonych zakresach mogą zostać opisane rozkładem jednorodnym.



Rys. 5.24. Jedno i dwuwymiarowe rozkłady gęstości prawdopodobieństwa dla parametrów modelowych 1009, 1011, 1014, 1031, 1033 uzyskane stosując metodę wstecznej kwantyfikacji niepewności opartej o analizę Bayesowską dla obliczeń wspólnych



Rys. 5.25. Jedno i dwuwymiarowe rozkłady gęstości prawdopodobieństwa dla parametrów modelowych 1009, 1011, 1014, 1031, 1033 uzyskane stosując metodę wstecznej kwantyfikacji niepewności opartej o analizę Bayesowską dla obliczeń oddzielnych

Pierwszym parametrem, dla którego obserwowane są znaczne różnice w uzyskanych dla obu podejść rozkładach jest parametr 1031. W przypadku obliczeń wspólnych obserwowane są dwa maksima, pierwsze dla wartości 1.5 oraz drugie dla wartości 6.27 która jest jednocześnie wartością odcięcia zakresu akceptowalnych wartości. Kształt tego rozkładu jest wynikiem opisanego powyżej kompromisu w uzyskiwaniu jednego rozkładu z trzech różnych rozkładów uzyskiwanych dla każdego z testów. W przypadku tego parametru różnice w lokalizacji maksimów były największe, co jest odzwierciedlone w ostatecznym sumarycznym rozkładzie parametru. W przypadku obliczeń wspólnych uzyskany zakres nie obejmuje wartości domyślnej (równiej 1) i jest ograniczony między wartościami 2 oraz 7.5. Rozkład nie posiada

żadnego charakterystycznego maksimum, jedynie obszar pomiędzy 2.75 a 4 gdzie częstości występowania wartości są największe, nie są one jednak znacznie większe niż częstości dla pozostałych wartości. Drugim parametrem, dla którego występują znacznie rozbieżności w kształcie wyznaczonych rozkładów jest parametr 1033. Zakres dla obliczeń wspólnych wynosi od 0.9 do 8.75, natomiast dla obliczeń oddzielnych od 1 do 6. W obliczeniach wspólnych zidentyfikowano jedno maksimum dla wartości pomiędzy 2 a 2.5, a następnie wartości od 3 do końca zakresu mogą być opisane rozkładem jednorodnym. W przypadku obliczeń oddzielnych zidentyfikowano dwa lokalne maksima, jednak rozkład parametru można przybliżyć za pomocą rozkładu jednorodnego.

5.2.4.2. Kwantyfikacja niepewności parametrów wejściowych stosując metodę opartą o uczenie maszyn

Kwantyfikację niepewności pięciu parametrów wejściowych dla zjawisk związanych z wymianą ciepłą w rdzeniu podczas fazy zalewania w trakcie awarii typu LBLOCA przeprowadzono również stosując opracowaną metodę opartą o uczenie maszyn, opisaną w [rozdziale 4.5.2](#).

Cechy dla obliczeń temperatury koszulki paliwowej oraz poziomu wody w rdzeniu – instalacja testowa Flecht Seaset

Cechy, czyli mierniki dla ilościowej analizy wyników obliczeń opracowano dla dwóch analizowanych wielkości wyjściowych. Dla temperatury koszulki paliwowej zidentyfikowano jedenaście cech. Pierwsze dwie cechy to maksymalna temperatura koszulki paliwowej w obliczeniach oraz ta wartość znormalizowana do wartości maksymalnej temperatury koszulki paliwowej w eksperymencie. Kolejne dwie cechy to czas, w którym dla obliczeń wyznaczono maksymalną temperaturę koszulki paliwowej oraz wartość ta znormalizowana do czasu, w którym zaobserwowano maksymalną temperaturę koszulki w eksperymencie. Kolejną cechą jest czas, w którym osiągnięto minimalną wartość temperatury (temperaturę nasycenia przy ciśnieniu w obudowie bezpieczeństwa) czyli czas schłodzenia najgorętszego pręta paliwowego wyznaczony w obliczeniach. Wartość ta znormalizowana względem odpowiedniego momentu czasowego zmierzonego w eksperymencie stanowi szóstą cechę. Kolejne dwie cechy wyznaczono korzystając z kwantyfikatorów dla ilościowej analizy wyników opisanych w metodzie CIRCE [89] zaproponowanej przez francuski IRSN. Siódmą cechą była więc całość obliczana do czasu osiągnięcia minimalnej wartości temperatury. Ósmą cechą był czas, w którym temperatura koszulki osiąga wartość równą temperaturze koszulki na

początku testu w obliczeniach. Dziewiątą cechą jest wartość cechy ósmej znormalizowana do odpowiadającego czasu w eksperymencie. Dwie pozostałe cechy są wyznaczone korzystając z równania 5.1. Są to odchylenie standardowe w przedziałach dziesięciosekundowych w całym czasie obliczeniowym oraz w okresie do czasu osiągnięcia minimalnej wartości temperatury.

Dla propagacji frontu zalewania również określono 11 cech, które opiszą cały przebieg zmian poziomu wody w rdzeniu. Pierwszą cechą jest czas zalania rdzenia w obliczeniach. Drugą cechą jest wartość cechy pierwszej znormalizowana względem odpowiadającej wartości eksperymentalnej. Trzecią cechą jest poziom wody w obliczeniach w czasie zalania rdzenia w eksperymencie. Czwartą cechą jest maksymalna wartość poziomu wody w rdzeniu. Natomiast kolejna cecha to normalizacja cechy numer cztery względem wysokości rdzenia (3.66 metra). Cecha szósta to całka obliczana do czasu zalania rdzenia. Siódmą cechą jest odchylenie standardowe obliczane w przedziałach dziesięciosekundowych zgodnie z równaniem 5.1. Ostatnie cztery cechy odpowiadają miernikom zastosowanym w globalnej analizie wrażliwości. Dwie z nich to wartość poziomu wody w rdzeniu w 1/3 czasu trwania eksperymentu oraz wartość poziomu wody w rdzeniu 2/3 czasu trwania eksperymentu. Dwie ostatnie to poprzednie wartości znormalizowane do odpowiadających im wartości zmierzonych w eksperymencie.

Klasy dla obliczeń temperatury koszulki paliwowej oraz poziomu wody w rdzeniu – instalacja testowa Flecht Seaset

Klasy stanowią dane wyjściowe z algorytmu uczenia maszyn jako reprezentacja stopnia zgodności między danymi pomiarowymi, a wynikami obliczeń. W przypadku wykonywania obliczeń dla instalacji Flecht Seaset i badania temperatury koszulki paliwowej zdefiniowano pięć klas o oznaczeniach od A – E:

- A. Satysfakcjonująca zgodność z danymi eksperymentalnymi w całym czasie trwania eksperymentu;
- B. Satysfakcjonująca zgodność w czasie do wystąpienia maksymalnej temperatury koszulki, jednak z nieodpowiednim wyznaczeniem czasu osiągnięcia temperatury minimalnej;
- C. Satysfakcjonująca zgodność w wyznaczaniu czasu osiągnięcia temperatury minimalnej, jednak maksymalna temperatura koszulki odbiegająca od danych pomiarowych;
- D. Brak zgodności z danymi eksperymentalnymi w całym czasie trwania eksperymentu;

- E. Wyniki znacznie odbiegające od przebiegu eksperymentalnego bądź obliczenia przerwane przez wystąpienia błędów.

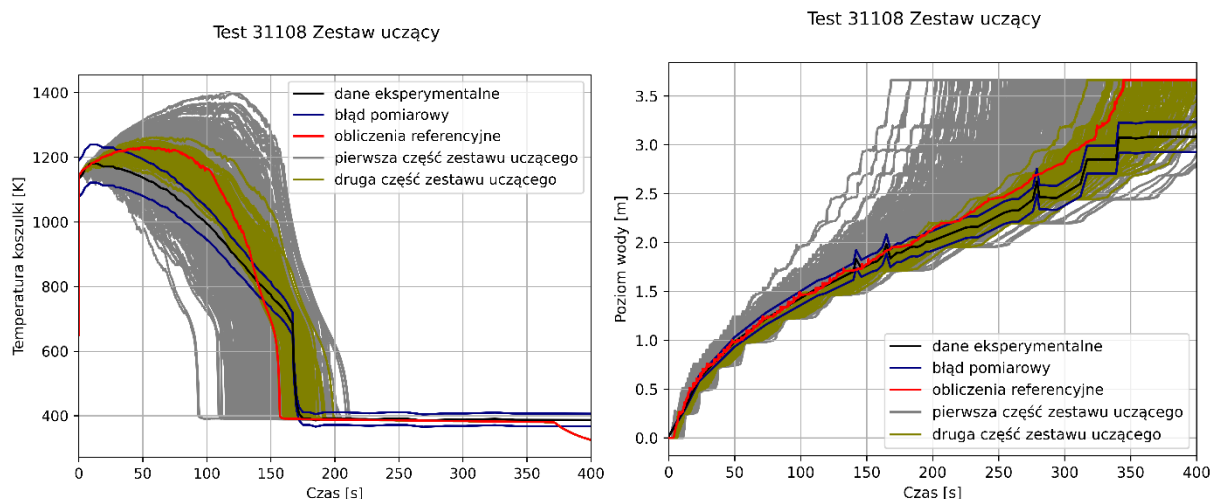
Dla poziomu wody w rdzeniu również zdefiniowano pięć klas:

- A. Satysfakcjonująca zgodność z danymi eksperymentalnymi w całym czasie trwania eksperymentu;
- B. Satysfakcjonująca zgodność w większości czasu przebiegu, jednak ostatecznie nieosiągnięcie całkowitego zalanania rdzenia;
- C. Zalanie rdzenia w czasie zbliżonym do eksperymentalnego, ale przebieg niedostatecznie zgodny z eksperymentalnym;
- D. Brak zgodności z danymi eksperymentalnymi w całym czasie trwania eksperymentu;
- E. Wyniki znacznie odbiegające od przebiegu eksperymentalnego bądź obliczenia przerwane przez wystąpienia błędów.

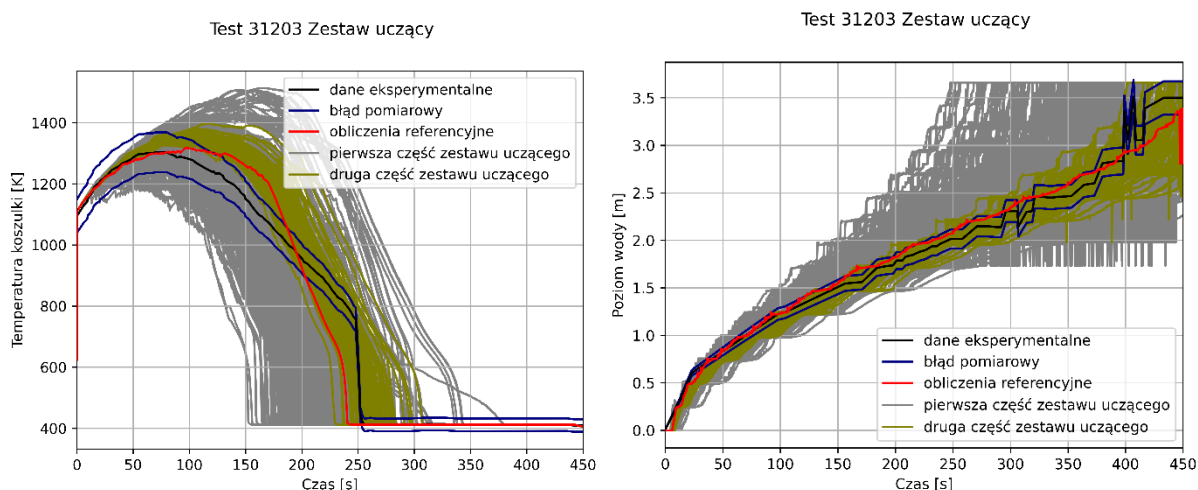
Zestaw uczący dla obliczeń temperatury koszulki paliwowej oraz poziomu wody w rdzeniu – instalacja testowa Flecht Seaset

Korzystając z doświadczeń z zastosowania metody kwantyfikacji niepewności opartej o uczenie maszyn dla instalacji Marviken zdecydowano się uwzględnić obliczenia dla trzech różnych testów w zbiorze uczącym. Zgodnie z przypisaniem przedstawionym w tabeli 5.5 do zestawu uczącego wybrano testy 31108, 31203 oraz 32114, posiadające różne warunki początkowe oraz przebiegi analizowanych wielkości wyjściowych. Dla wszystkich testów wykonano początkowo 420 obliczeń, które stanowiły pierwszą część zestawu uczącego. W obliczeniach tych wartości parametrów losowano z rozkładu jednorodnego ograniczonego zakresem od 0 do 6 dla wszystkich parametrów poza 1011. Dla tego parametru, po analizie wyników z analizy wrażliwości, dolny zakres ograniczono do wartości 0.5. Obliczenia z pierwszej części zestawu przedstawiono kolorem szarym dla testu 31108 na rysunku 5.26, dla testu 31203 na rysunku 5.27 oraz testu 32114 na rysunku 5.28. Szeroki zakres zmienności parametrów zgodnie z założeniami spowodował, że wiele obliczeń zostało przypisanych do klasy D. Na podstawie analizy wyników wyznaczono zakresy parametrów, w których większość wyników znalazła się w klasach A, B oraz C. Na tej podstawie wykonano dodatkowe 100 obliczeń stanowiące część drugą zestawu uczącego. Wartości parametrów losowano z rozkładów jednorodnych o następujących zakresach: dla 1009 od 0.4 do 3.5, dla 1011 od 1.5 do 2.5, dla 1014 od 0.25 do 1.5, dla 1031 od 2.5 do 4.25 a dla 1033 od 3.5 do 4.5. Obliczenia

z drugiej części zestawu przedstawiono kolorem oliwkowym dla testu 31108 na rysunku 5.26, dla testu 31203 na rysunku 5.27 oraz testu 32114 na rysunku 5.28. Obliczenia te pozwoliły na przypisanie dużej liczby obliczeń do klas A, B oraz C dla testów 31108 oraz 31203, natomiast dla testu 32114 szczególnie dla poziomej wody w rdzeniu większość obliczeń została zaklasyfikowana do klasy D.



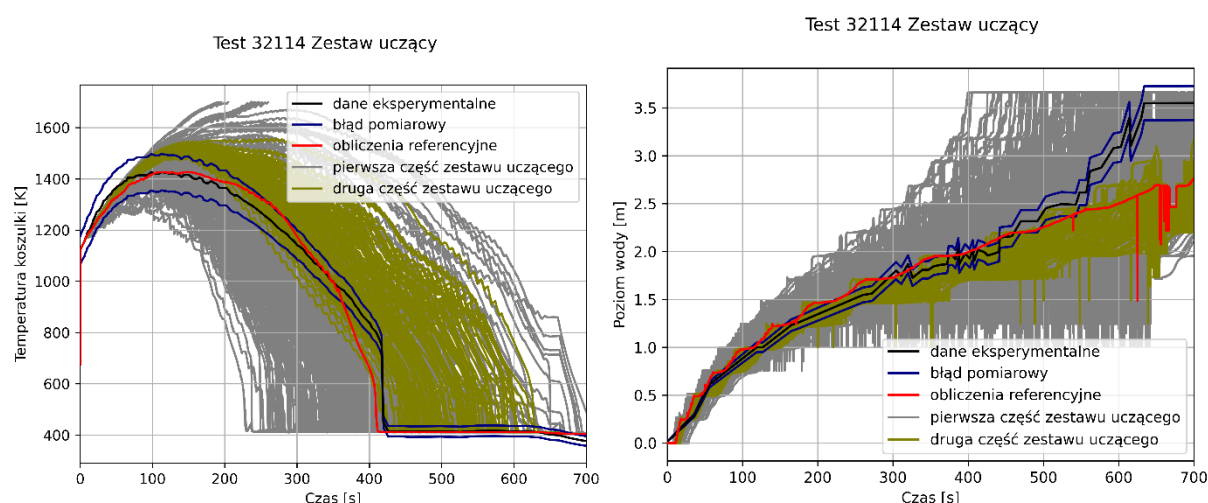
Rys. 5.26. Zestaw uczący dla obliczeń temperatury koszulki paliwowej (po lewej) oraz poziomu wody w rdzeniu (po prawej) dla testu 31108 w instalacji Flecht Seaset w podziale na pierwszą część zestawu (kolor szary), oraz drugą część zestawu (kolor oliwkowy)



Rys. 5.27. Zestaw uczący dla obliczeń temperatury koszulki paliwowej (po lewej) oraz poziomu wody w rdzeniu (po prawej) dla testu 31203 w instalacji Flecht Seaset w podziale na pierwszą część zestawu (kolor szary), oraz drugą część zestawu (kolor oliwkowy)

Dla wszystkich obliczeń wykonanych dla każdego z testów w zbiorze uczącym dokonano porównania z danymi eksperymentalnymi i na tej podstawie manualnie dokonano przypisania do jednej z pięciu zdefiniowanych powyżej klas. Wyniki przypisania obliczeń do klas

przedstawiono w tabeli 5.8. Najwięcej obliczeń przypisano do klasy D, łącznie 39% w przypadku temperatury koszulki paliwowej oraz 53% dla propagacji frontu zalewania. Do klas B oraz C przypisano ponad 20% obliczeń dla przebiegu zmian temperatury oraz 16% dla poziomu wody w rdzeniu. Klasa A stanowiła 8% wszystkich wykonanych obliczeń w analizie temperatury koszulki paliwowej. Najwięcej obliczeń do klasy A przypisano do testu 31108, a najmniej dla testu 32114, gdzie obliczenia referencyjne były zbieżne z danymi eksperymentalnymi. W przypadku poziomu wody w rdzeniu ponownie najwięcej obliczeń dla klasy A przypisano dla testu 32108 natomiast najmniej dla 31203. Sumarycznie 88 obliczeń więcej zostało przypisanych do klasy A dla poziomu wody w rdzeniu, co ostatecznie daje 13.6% wszystkich obliczeń.



Rys. 5.28. Zestaw uczący dla obliczeń temperatury koszulki paliwowej (po lewej) oraz poziomu wody w rdzeniu (po prawej) dla testu 32114 w instalacji Flecht Seaset w podziale na pierwszą część zestawu (kolor szary), oraz drugą część zestawu (kolor oliwkowy)

Tabela 5.8. Liczba obliczeń przypisanych manualnie do pięciu klas dla testów tworzących zbiór uczący dla instalacji Flecht Seaset

		Klasa A	Klasa B	Klasa C	Klasa D	Klasa E
Temperatura koszulki paliwowej	Test 31108	64	167	117	172	0
	Test 31203	34	165	137	184	0
	Test 32114	26	115	117	251	11
	Suma	124	447	371	607	11
	Udział [%]	7.9%	28.7%	23.8%	38.9%	0.7%
Propagacja frontu zalewania	Test 31108	111	36	87	286	0
	Test 31203	45	127	115	233	0
	Test 32114	56	94	52	307	11
	Suma	212	257	254	826	11
	Udział [%]	13.6%	16.5%	16.3%	52.9%	0.7%

Algorytm nadzorowanego uczenia maszyn zostanie zastosowany oddzielnie dla obu analizowanych wielkości wyjściowych. Stąd dla wszystkich 1560 obliczeń z zestawu uczącego dokonano obliczeń wartości cech opisanych powyżej wpierw dla koszulki paliwowej a następnie dla propagacji frontu zalewania. Następnie stosując klasyfikator lasów losowych (opisany w [rozdziale 4.5.2.4](#)) algorytm uczenia maszyn nauczył się relacji między cechami a klasami.

Zestaw testowy dla obliczeń temperatury koszulki paliwowej oraz poziomu wody w rdzeniu – instalacja testowa Flecht Seaset

Zestaw testowy składa się z obliczeń wykonanych dla trzech testów 30817, 31302 oraz 31504. Dla każdego testu wykonano 24 576 obliczeń, stosując do wygenerowania wartości parametrów generator losowy Saltelliego. Liczba wykonanych obliczeń jest w tej metodzie większa niż w metodzie Bayesowskiej, ponieważ możliwe jest prowadzenie wielu obliczeń równoległe. Wszystkie pliki wejściowe dla kodu TRACE zostają przygotowane jednocześnie po wylosowaniu wartości parametrów. W metodzie Bayesowskiej nie można wykonywać dla jednego testu więcej niż jednego obliczenia, ponieważ plik wsadowy dla następnego kroku jest tworzony dopiero po wylosowaniu wartości parametrów z rozkładów propozycji, co jest uzależnione od obecnego kroku. Należy zaznaczyć, że w metodzie uczenia maszyn wylosowany raz zestaw parametrów jest stosowany do opracowania plików wejściowych dla wszystkich testów w zestawie uczącym. Tak, więc przykładowo obliczenia numer 10 posiadają taki sam zestaw parametrów dla wszystkich trzech testów.

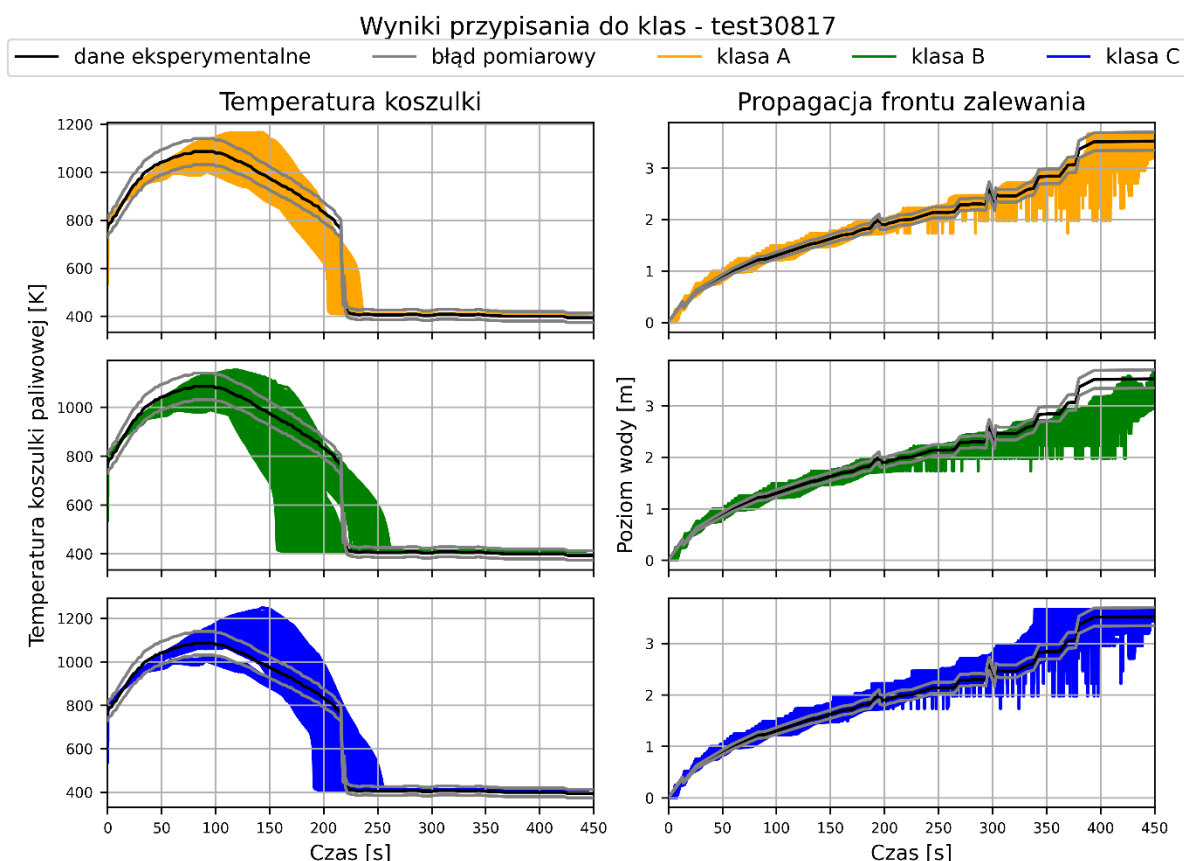
Zakresy, z których losowano wartości pięciu parametrów określono na podstawie analizy wyników uzyskanych dla zestawu uczącego. Wartości parametrów losowano z rozkładów jednorodnych ograniczonych zakresami: dla 1009 od 0.5 do 3.75, dla 1011 od 0.75 do 3, dla 1014 od 0.25 do 3, dla 1031 od 1 do 7, dla 1033 od 1 do 7.

Po wykonaniu obliczeń dla wszystkich testów ze zbioru testowego, dla każdego z tych obliczeń wyznaczone zostały wartości cech dla obu badanych wielkości wyjściowych. Zgodnie z opisem w poprzedniej sekcji model uczenia maszyn zna również relację między cechami a klasami dla poszczególnej wielkości wyjściowej. Na tej podstawie algorytm przypisywał jedną z opisanych powyżej pięciu klas dla każdego z wykonanych obliczeń oddzielnie dla temperatury koszulki paliwowej oraz dla poziomu wody w rdzeniu. W tabeli 5.9 przedstawiono liczbę obliczeń przypisanych przez model uczenia maszyn dla każdego z testów wykonanych w instalacji

Flecht Seaset i tworzących zbiór testowy, w podziale na dwie analizowane wielkości wyjściowe.

Tabela 5.9. Liczba obliczeń przypisanych przez algorytm uczenia maszyn do pięciu klas dla testów tworzących zbiór testowy dla instalacji Flecht Seaset

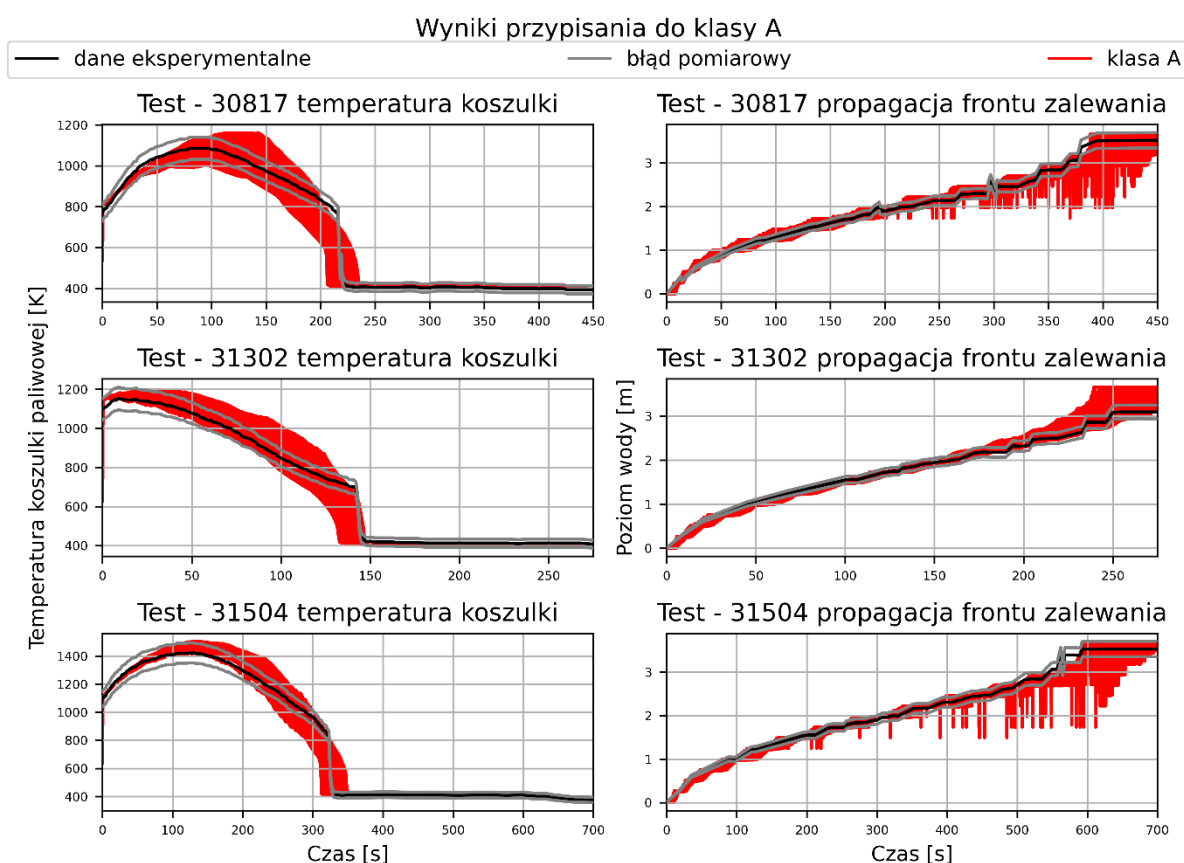
		Klasa A	Klasa B	Klasa C	Klasa D	Klasa E
Temperatura koszulki paliwowej	Test 30817	6718	10921	5057	1880	0
	Test 31302	1388	8853	6206	8129	0
	Test 31504	2406	651	8449	13052	18
	Suma	10512	20425	19712	23061	18
	Udział [%]	14.3%	27.7%	26.7%	31.3%	0.0%
Propagacja frontu zalewania	Test 30817	5126	4040	11115	4295	0
	Test 31302	2400	5951	5738	10487	0
	Test 31504	6698	8912	4165	4767	34
	Suma	14224	18903	21018	19549	34
	Udział [%]	19.3%	25.6%	28.5%	26.5%	0.0%



Rys. 5.29. Przypisanie do klas A, B, C wykonane z zastosowaniem metody uczenia maszyn dla testu 30817 w instalacji Flecht Seaset

Zarówno dla temperatury koszulki paliwowej jak i propagacji frontu zalewania większość wyników przypisane zostało do klas B, C oraz D. Udział klasy D nie jest aż tak dominujący jak w obliczeniach dla instalacji testowej Marviken. Do klasy A dla temperatury koszulki

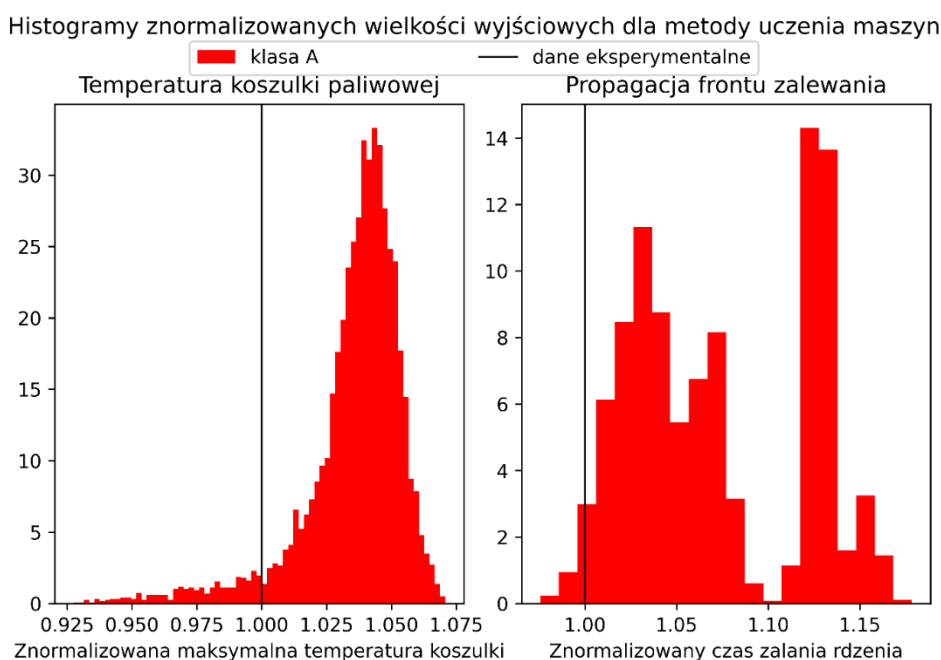
paliwowej przypisano 14% obliczeń, natomiast dla poziomu wody w rdzeniu aż 19%. W obu przypadkach najmniej obliczeń algorytm przypisał do klasy A dla testu 31302. Przypisanie do klas A, B i C dla testu 30817 przedstawiono na rysunku 5.29. Wyniki przypisania do wszystkich pięciu klas dla wszystkich testów przedstawiono w [aneksie](#). W przypadku zaprezentowanych wyników dla testu 30817 uzyskane wyniki dla klasy A pokrywają dane eksperymentalne i błąd pomiarowy, jednak rozrzut wyników nie jest znaczący. Podobnie jak w przypadku obliczeń z zastosowaniem metody Bayesowskiej, tak i przy zastosowaniu uczenia maszyn algorytm do klasy A przypisywał obliczenia w których istotny był moment osiągnięcia temperatury minimalnej. W przypadku poziomu wody w rdzeniu można zaobserwować, że czas zalania rdzenia dla obliczeń w klasie A był bardzo zbliżony do obliczeń eksperymentalnych. Obserwowane są oscylacje, jednak ogólnie obliczenia w pełni pokrywają dane eksperymentalne.



Rys. 5.30. Przypisanie do klas A dla obu wielkości dla testów z zestawu testowego dla obliczeń wykonanych z zastosowaniem metody uczenia maszyn w instalacji Flecht Seaset

Podobnie jak dla obliczeń oddzielnych wykonanych stosując metodę Bayesowską tak i w tym wypadku konieczne jest odnalezienie obliczeń, które dają satysfakcjonujące wyniki dla temperatury koszulki paliwowej oraz poziomu wody w rdzeniu. W obrębie wyników dla

każdego testu poszukiwano więc takich obliczeń dla których algorytm przypisał klasę A dla obu wielkości. Dla testu 30817 zidentyfikowano 4127 takich obliczeń, 981 dla testu 31302 oraz 1650 dla testu 31504. Obliczenia te przedstawiono na rysunku 5.30. Dla każdego z testów odnalezienie obliczeń przypisanych do dwóch klas A pozwoliło na prawidłowe odwzorowanie danych eksperymentalnych. Najbardziej dokładne odwzorowanie zaobserwowano dla testu 31504, gdzie wyniki obliczeń dla temperatury koszulki paliwowej jedynie nieznacznie wykraczają poza określony błąd pomiarowy. W przypadku poziomu wody w rdzeniu najlepsze dopasowanie obserwowane jest dla testu 31302. Ilościową miarą uzyskanych wyników było porównanie maksymalnej temperatury koszulki paliwowej oraz czasu zalania rdzenia z danymi doświadczalnymi. Histogramy uzyskanych znormalizowanych wartości przedstawiono na rysunku 5.31. Rozbieżności maksymalnej temperatury koszulki paliwowej wyników w podwójnej klasie A nie przekroczyły 7% w odniesieniu do danych eksperymentalnych. W przypadku czasu zalania rdzenia rozbieżności były większe, sięgały maksymalnie 15%. W przypadku poziomu wody jest to akceptowalne, ponieważ dla testu 31302 dla danych eksperymentalnych maksymalny poziom wody wynosił 3.1 metra (w dokumentacji nie podano uzasadnienia dla takiej wartości). Stąd nawet przy zbieżnym przebiegu wartości osiągnięcie maksymalnej wysokości rdzenia, stałej w modelu obliczeniowym, zajmuje więcej czasu. Obliczenia zaklasyfikowane do podwójnej klasy A stanowią dobrą podstawę do wyznaczenia na ich podstawie zakresu zmienności parametrów modelowych.

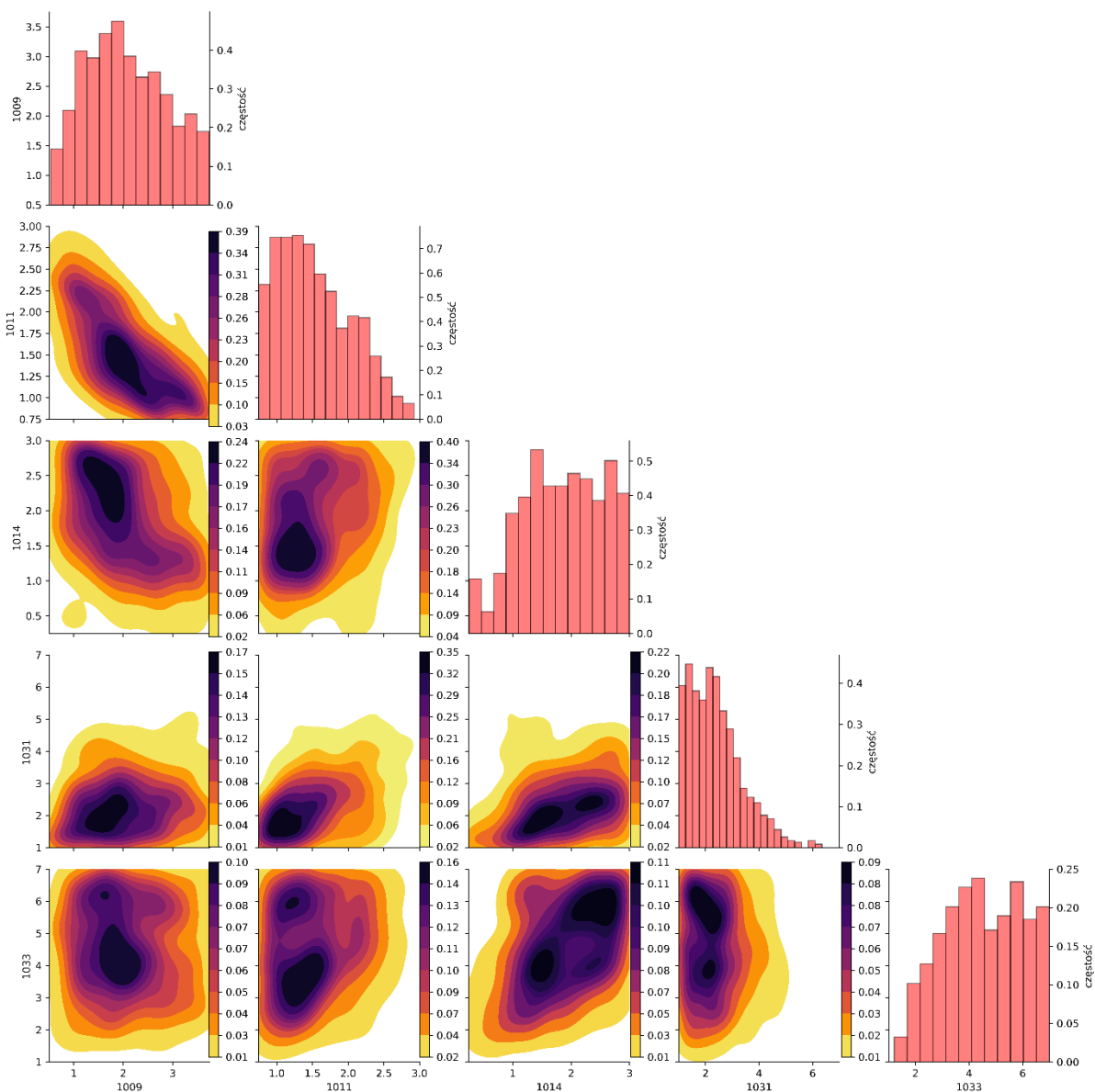


Rys. 5.31. Znormalizowana maksymalna temperatura koszulki paliwowej oraz maksymalny poziom wody dla wyników zaklasyfikowanych do klasy A dla obu wielkości

Określenie rozkładów gęstości prawdopodobieństwa parametrów

Obliczenia zaklasyfikowane do klasy A dla obu wielkości wyjściowych dla każdego testu stanowiły podstawę do określenia wartości parametrów, z których utworzone zostaną ich rozkłady. Idealnym rozwiązaniem byłoby znalezienie takiego zestawu parametrów dla którego dla wszystkich trzech testów przypisano podwójną klasę A. W całym zbiorze 24 576 obliczeń nie było jednak ani jednego takiego przypadku. Zebrano, więc wszystkie zestawy parametrów, dla których przynajmniej w jednym teście przypisano obu wielkościom wyjściowym klasę A, a następnie odrzucono te zestawy parametrów, dla których w choć jednym z pozostałych czterech przypadków przypisano klasę D bądź E. W ten sposób ostateczny zbiór ograniczono do 897 zestawów parametrów. Na podstawie tych wartości określono rozkłady gęstości prawdopodobieństwa pięciu badanych parametrów wewnętrznych kodu, które przedstawiono na rysunku 5.32.

Rozkład parametru 1009 charakteryzuje się obecnością maksimum w okolicach wartości 1.9. Najmniej prawdopodobne wartości znajdują się w zakresie 0.5 – 1.0 oraz od 3.0 do 3.75. Dla parametru 1011 75% wartości znajduje się w przedziale od 0.75 do 2, pozostałe 25% w przedziale od 2 do 3. Obszar o największym prawdopodobieństwie uzyskania poprawnych wyników obejmuje zakres między 1 a 1.5. Dla parametru 1014 większość prawdopodobnych wartości jest w zakresie 1.15 – 3.0. Wartości poniżej 1.15 są bardzo mało prawdopodobne. Dla parametru 1031 75% wartości znajduje się w przedziale od 1 do 3. Pozostałe 25% w przedziale od 3 do 6.5, z czego wartości powyżej 5 są bardzo mało prawdopodobne. Jest to pierwszy parametr, dla którego wyznaczona górna wartość zakresu prawdopodobieństwa (równa 6.35) jest mniejsza niż górna wartość zakresu, z którego losowano parametry (równa 7). Najbardziej prawdopodobne wartości dla parametru 1033 znajdują się w przedziale od 3 do 7. Nie można jednak zidentyfikować jednego wyraźnie zarysowanego maksimum. Najmniejsza możliwa wartość to 1.2, która jest większa od wartości pierwotnie ograniczającej zakres losowania wartości parametrów (równa 1).

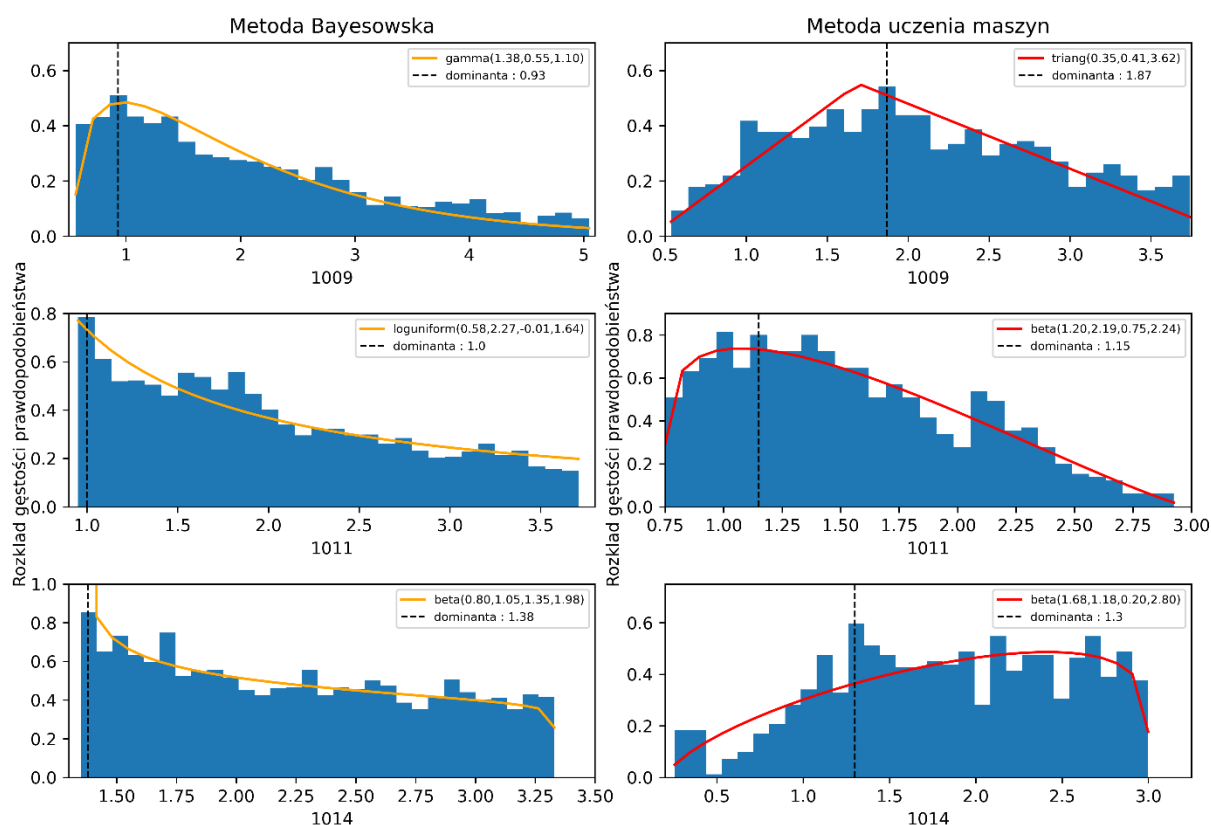


Rys. 5.32. Jedno i dwuwymiarowe rozkłady gęstości prawdopodobieństwa dla parametrów modelowych 1009, 1011, 1014, 1031, 1033 uzyskane stosując metodę wstecznej kwantyfikacji niepewności opartą o uczenie maszyn

5.2.4.3. Porównanie wyników oraz propozycja rozkładów gęstości prawdopodobieństwa

Do końcowej analizy wyników rozkładów gęstości prawdopodobieństwa dla metody Bayesowskiej wybrano rozkłady uzyskane w obliczeniach oddzielnych. Rozkłady tworzone były na podstawie 3057 zestawów parametrów, tymczasem dla obliczeń wspólnych było ich tylko 1574. Dodatkowo zaobserwowano większą zgodność między zestawami parametrów uzyskanymi dla trzech testów, dla których wykonywano obliczenia. Pozwoliły one również na uzyskanie rozkładów które łatwiej aproksymować funkcjami ciągłymi. Dopasowanie funkcjami ciągłymi dla uzyskanych rozkładów gęstości prawdopodobieństwa dla obu metod

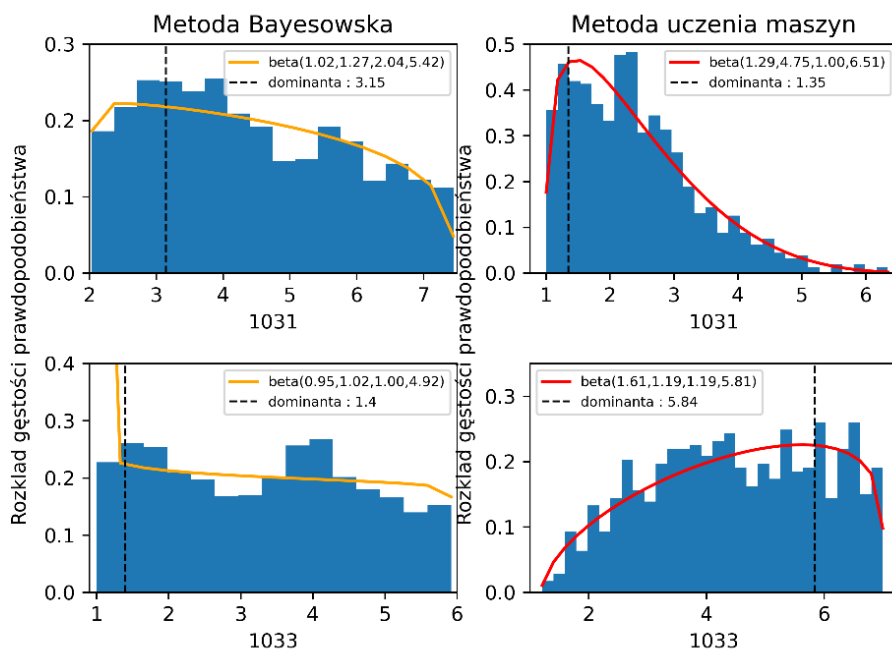
kwantyfikacji niepewności przedstawiono na rysunkach 5.33 oraz 5.34. Dla parametru 1009 zaproponowano rozkład gamma dla metody Bayesowskiej oraz rozkład trójkątny dla metody uczenia maszyn. Obserwowalne są różnice między uzyskanymi rozkładami, jednak w obu przypadkach większość najbardziej prawdopodobnych wartości jest w przedziale od 0.9 do 3. Dla parametru 1011 zaproponowano rozkład logarytmiczno-jednorodny dla metody Bayesowskiej oraz beta dla metody uczenia maszyn. Kształt obu rozkładów jest jednak bardzo podobny, poza wartościami poniżej 0.95 które są odcięte dla metody Bayesowskiej. Wartości najlepszego szacowania również są zbliżone i wynoszą 1 oraz 1.15. Dla parametru 1014 różnice występują dla wartości poniżej 1.35 które są odcięte w metodzie Bayesowskiej. Natomiast dalsze rozkłady wartości są podobne i odcinają się w okolicach wartości równej 3. Dla obu rozkładów zaproponowano ciągły rozkład beta.



Rys. 5.33. Propozycje ciągłych rozkładów gęstości prawdopodobieństwa dla parametrów 1009, 1011 oraz 1014 w wyniku zastosowania dwóch metod wstecznej kwantyfikacji niepewności

Dla parametru 1031 zauważalne są znaczne różnice w uzyskanych rozkładach oraz zakresach ich zmienności. Podobnie znaczne różnice można zaobserwować w wartościach najlepszego szacowania. Jedyna część wspólna to zakres między 2 a 3.25 gdzie wartości dają wysokie prawdopodobieństwa. Dla parametru 1033 podobieństwa występują w zakresie 2.5 a 6. Różnice

występują w wartościach poniżej 2.5, jak również w wartościach najlepszego szacowania. Przedstawione rozkłady gęstości prawdopodobieństwa parametrów będą stanowiły podstawę do losowaniach z nich wartości parametrów w obliczeniach sprawdzających i walidacyjnych. Pozwoli to na określenie, które rozkłady pozwalają na uzyskanie rezultatów bliższych do danych eksperymentalnych.

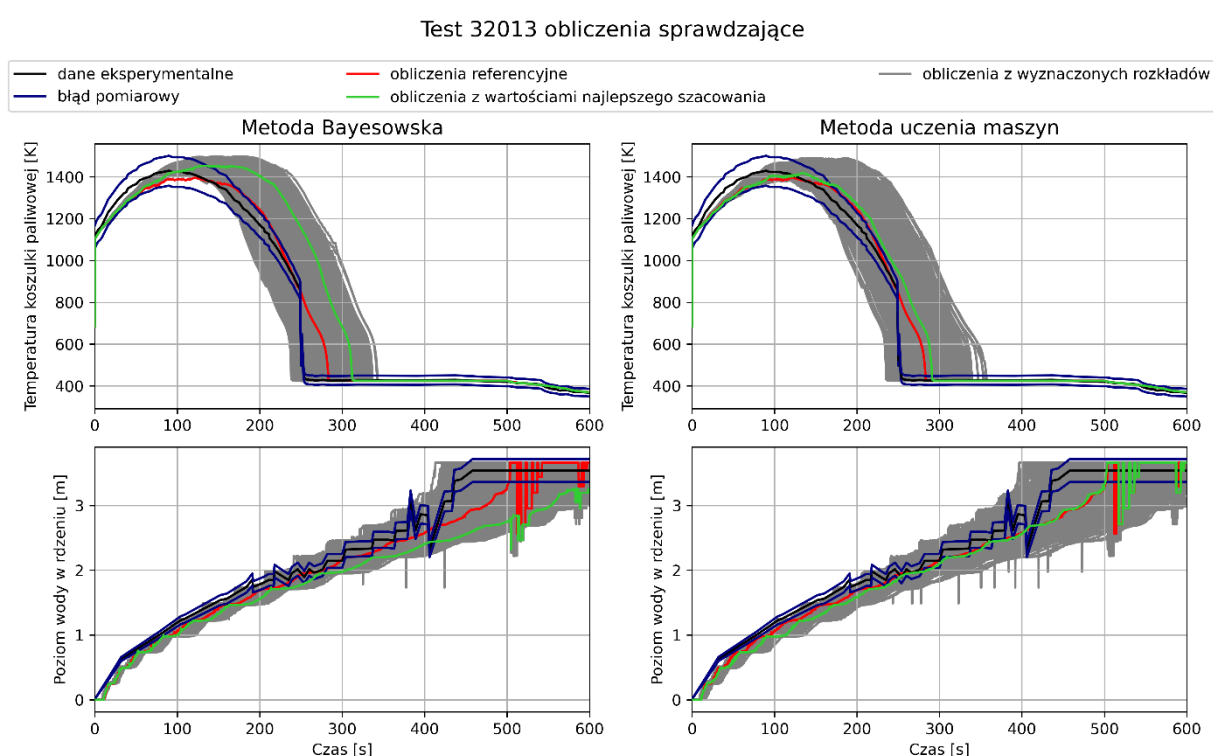


Rys. 5.34. Propozycje ciągłych rozkładów gęstości prawdopodobieństwa dla parametrów 1031 oraz 1033 w wyniku zastosowania dwóch metod wstecznej kwantyfikacji niepewności

5.2.5. Sprawdzenie poprawności uzyskanych rozkładów gęstości prawdopodobieństwa (krok 13)

Sprawdzenie czy wartości losowane z uzyskanych rozkładów gęstości prawdopodobieństwa pozwolą na uzyskanie wyników zbliżonych z danymi eksperymentalnymi przeprowadzono dla czterech testów oznaczonych etykietą „walidacja” w tabeli 5.5. Wykonano po 1000 obliczeń dla każdego z tych czterech testów stosując rozkłady uzyskane wykorzystując metodę Bayesowską oraz uczenia maszyn, opisane w poprzedniej sekcji. Losowania parametrów dokonano stosując proste próbkowanie losowe, a przedstawione rezultaty przedstawiają 95 percentyl zgodnie z formułą Wilksa. Dodatkowo dla każdego testu wykonano obliczenia z wartościami najlepszego szacowania parametrów. Wyniki dla trzech testów (po jednym z każdej grupy podobieństwa przedstawionych w tabeli 5.5) przedstawiono na rysunkach 5.35, 5.36 oraz 5.37.

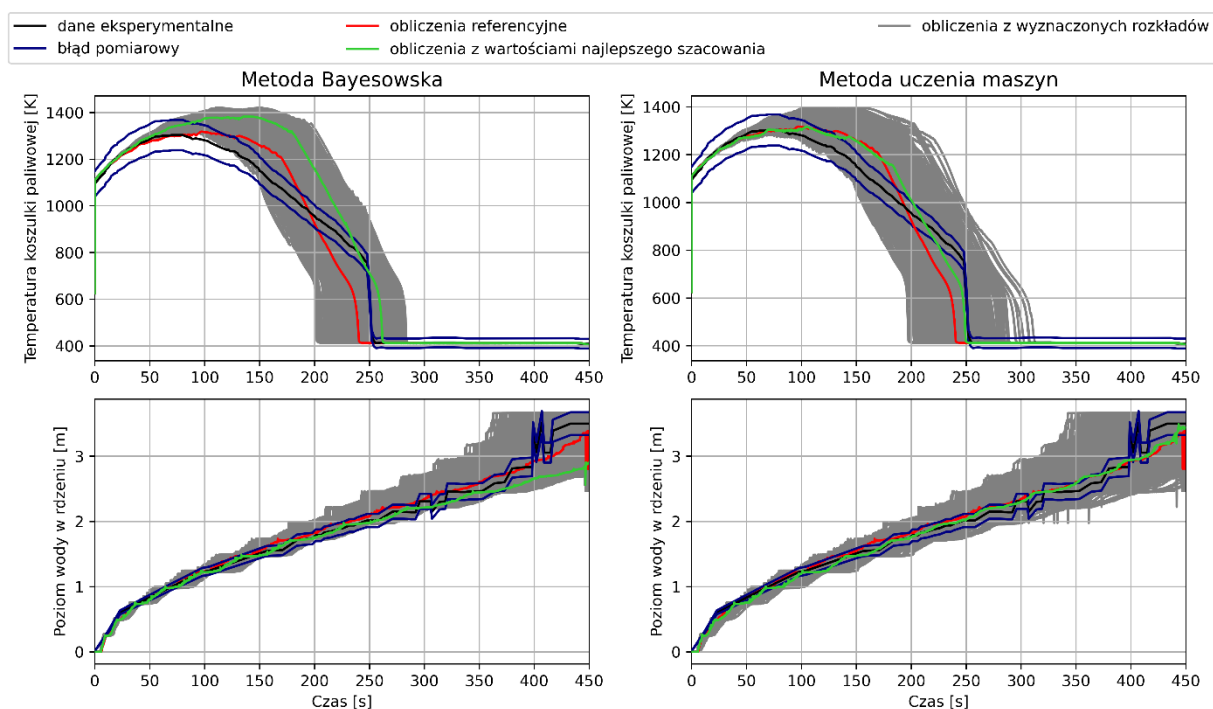
Obliczenia walidacyjne wykonano dla dwóch testów z grupy pierwszej. Lepsze rezultaty uzyskano dla testu 32013 przedstawionego na rysunku 5.35. Dla obu metod uzyskane zakresy niepewności są podobne i obejmują dane eksperymentalne. Różnice w wartości maksymalnej temperatury koszulki nie są większe niż 5%, jednak same zakresy rozbieżności dla temperatury koszulki są dość szerokie. Dla poziomu wody przedziały nie są bardzo szerokie, choć można zaobserwować, że w zbiorze walidacyjnym wiele jest wyników, dla których czas zalewania rdzenia był przeszacowany. Obliczenia z wartościami najlepszego szacowania w obu przypadkach nie pozwalają na poprawę zbieżności z danymi eksperymentalnymi w porównaniu z obliczeniami referencyjnymi.



Rys. 5.35. Zakresy niepewności temperatury koszulki paliwowej oraz poziomu wody w rdzeniu utworzone przez wyniki obliczeń sprawdzających dla testu 32013

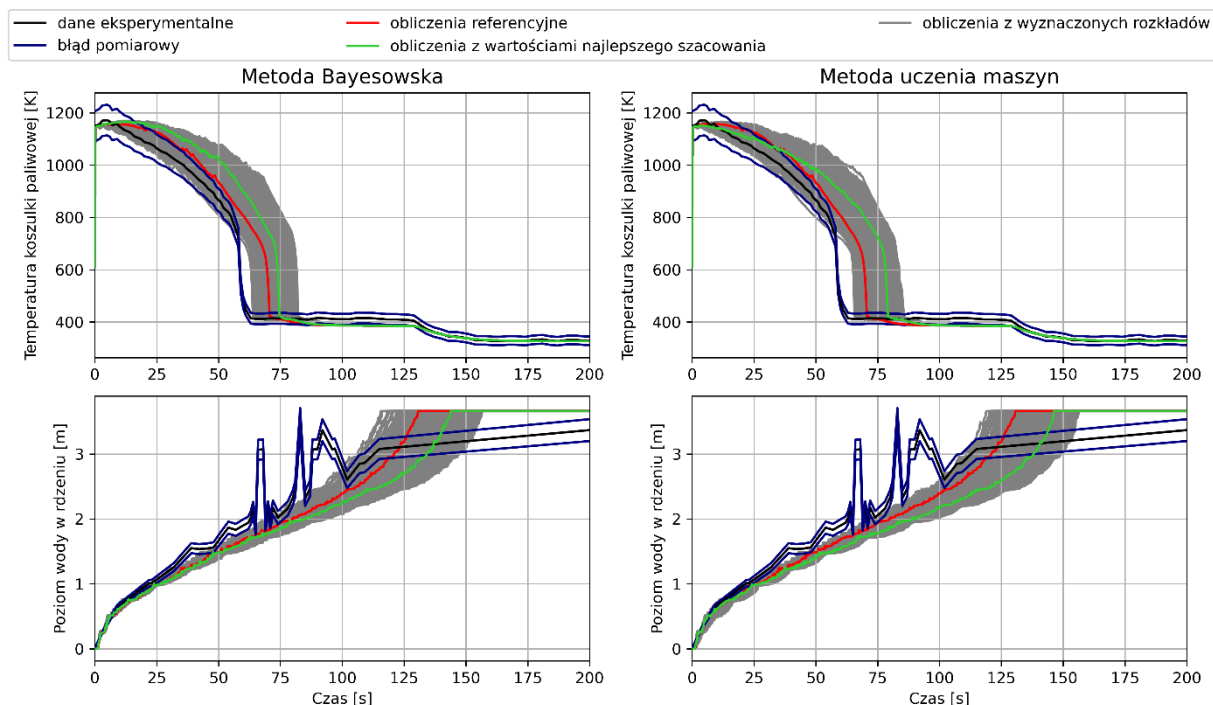
Obliczenia dla testu 31203 z grupy drugiej przedstawiono na rysunku 5.36. Obliczenia z wartościami najlepszego szacowania pozwalają na uzyskanie odpowiedniego dopasowania do danych eksperymentalnych dla metody uczenia maszyn. Zakresy rozbieżności wyników są węższe dla metody Bayesowskiej zarówno dla temperatury koszulki paliwowej jak również poziomu wody w rdzeniu. Stosując obie metody uzyskuje się jednak wyniki obejmujące dane eksperymentalne, choć w większości wypadków przeszacowane są zarówno maksymalne temperatury koszulki paliwowej jak i czas, w którym to następuje. Niemniej jednak wyniki są satysfakcjonujące dla obu zastosowanych metod.

Test 31203 obliczenia sprawdzające



Rys. 5.36. Zakresy niepewności temperatury koszulki paliwowej oraz poziomu wody w rdzeniu utworzone przez wyniki obliczeń sprawdzających dla testu 31203

Test 31701 obliczenia sprawdzające



Rys. 5.37. Zakresy niepewności temperatury koszulki paliwowej oraz poziomu wody w rdzeniu utworzone przez wyniki obliczeń sprawdzających dla testu 32013

Wyniki obliczeń dla testu 31701 (grupa trzecia) przedstawiono na rysunku 5.37. Obliczenia nie pozwoliły na poprawne odwzorowanie czasu, w którym temperatura dla gorącego pręta paliwowego osiąga temperaturę minimalną. W każdym wypadku są przeszacowane, podobnie jak całkowite przebiegi temperatury koszulki w obliczeniach z wartościami najlepszego szacowania. Przebieg eksperymentalny poziomu wody w rdzeniu obarczony jest widocznymi oscylacjami między 50 a 100 sekundą co było niemożliwe do uchwycenia w obliczeniach. Czas zalania rdzenia następuje szybciej w obliczeniach, co wydaje się zgodne z przebiegiem temperatury koszulki, dane eksperymentalne są w tym wypadku wątpliwej jakości.

Przedstawione wyniki obliczeń sprawdzających pokazują, że uzyskane rozkłady pozwalały na uzyskanie względnie szerokich rozkładów, jednakże w większości wypadków pokrywających dane eksperymentalne. Należy pamiętać, że uzyskane rozkłady są kompromisem dla wyników obliczeniowych uzyskiwanych dla testów z trzech grup o różnych warunkach początkowych. Trudności w określeniu takich rozkładów gęstości prawdopodobieństwa, które pozwalałyby na uzyskanie bardzo dobrych wyników w każdych warunkach są, więc zrozumiałe. Podejście, w którym poszukuje się jednego zestawu parametrów dla różnorodnych warunków skutkuje uzyskiwaniem szerszych zakresów niepewności badanych wielkości wyjściowych. Rozkłady te są jednak bardziej uniwersalne. Aby uzyskać dokładniejsze dopasowanie do wyników pomiarowych oraz węższe zakresy niepewności należałoby opracować trzy zestawy rozkładów gęstości parametrów, oddzielne dla każdej z grup podobieństwa. Rozkłady te byłyby jednak mało uniwersalne i dedykowane do stosowania tylko w ściśle określonych warunkach początkowych i brzegowych. Dodatkowo uzyskane rezultaty obliczeń sprawdzających po raz kolejny wskazują jak istotne są obliczenia referencyjne. Obliczenia te de facto mają kluczowy wpływ na wyznaczenie rozkładów a potem również na wyniki obliczeń sprawdzających. W obliczeniach referencyjnych dla testów stanowiących zbiór testowy (oraz podstawę dla obliczeń MCMC) czas osiągnięcia minimalnej temperatury koszulki był zawsze niedoszacowany. Tymczasem dla dwóch spośród testów walidacyjnych (32013 oraz 31701) czasy te były przeszacowane, co skutkowało przesunięciem uzyskanych rozkładów w kierunku przeszacowanych obliczeń referencyjnych. Gdyby dla testu 31701 obliczenia referencyjne nie doszacowały czasu osiągnięcia temp minimalnej to można zakładać, że uzyskane rezultaty z obliczeń sprawdzających pozwoliłyby na pokrycie całego przebiegu eksperymentalnego. Należy również pamiętać, że uzyskane parametry muszą pozwalać na uzyskanie akceptowalnych wyników nie dla jednej wielkości wyjściowej, lecz dwóch, dla których najlepsze wartości parametrów często były rozbieżne. Zaobserwowano również, że o ile

konkretne zestawy parametrów z danego zakresu pozwalały na osiągnięcie odpowiedniego dopasowania, to już próbkowanie losowe z tych zakresów nie daje aż tak wąskich przedziałów dopasowania. Oznacza to, że w obrębie wyznaczonych rozkładów dla pięciu parametrów istnieją zarówno kombinacje pozwalające na uzyskanie wyników zbieżnych z eksperymentem, jak i kombinacje, dla których wyniki są znacznie gorzej dopasowane do danych eksperymentalnych.

6. Walidacja uzyskanych rozkładów gęstości prawdopodobieństwa parametrów – elementy piąty i szóstki metodyki

Walidacja uzyskanych rozkładów gęstości prawdopodobieństwa parametrów wejściowych wykonywana jest zgodnie z elementami piątym oraz szóstym opracowanej metodyki. W elemencie piątym wykonywane są obliczenia walidacyjne dla instalacji eksperymentalnej, aby była możliwość porównania wyników obliczeń z danymi pomiarowymi. Zgodnie z opisem w rozdziale [czwartym](#) instalacją eksperymentalną dla której wykonane zostaną obliczenia walidacyjne jest instalacja LOFT. W instalacji tej wykonywano między innymi doświadczenia odwzorowujące awarię typu LBLOCA w elektrowni jądrowej, co pozwoliło na zebranie danych doświadczalnych dla zjawisk wymiany ciepła w rdzeniu oraz wypływu przez rozerwanie. Aby możliwe było wykonanie obliczeń walidacyjnych, należało opracować model matematyczny instalacji oraz wykonać obliczenia stanu ustalonego oraz referencyjne obliczenia stanu nieustalonego, co przedstawiono w rozdziałach 6.1.1 oraz 6.1.2. Obliczenia walidacyjne zaprezentowane są w rozdziale 6.1.3.

W elemencie szóstym metodyki wykonywane są obliczenia niepewności dla modelu reaktora energetycznego oraz wybranego scenariusza awaryjnego, czyli awarii typu LBLOCA. Obliczenia wykonane zostały dla trzypętłowego reaktora typu PWR o mocy elektrycznej 900 MW. W przypadku obliczeń dla reaktora energetycznego brak jest oczywiście danych eksperymentalnych oraz obliczeń referencyjnych dla elektrowni. Wyniki obliczeń walidacyjnych porównywane są więc z kryteriami akceptacji określonymi dla scenariusza awarii typu LBLOCA. Model reaktora PWR oraz wyniki obliczeń niepewności przedstawiono w rozdziale 6.2.

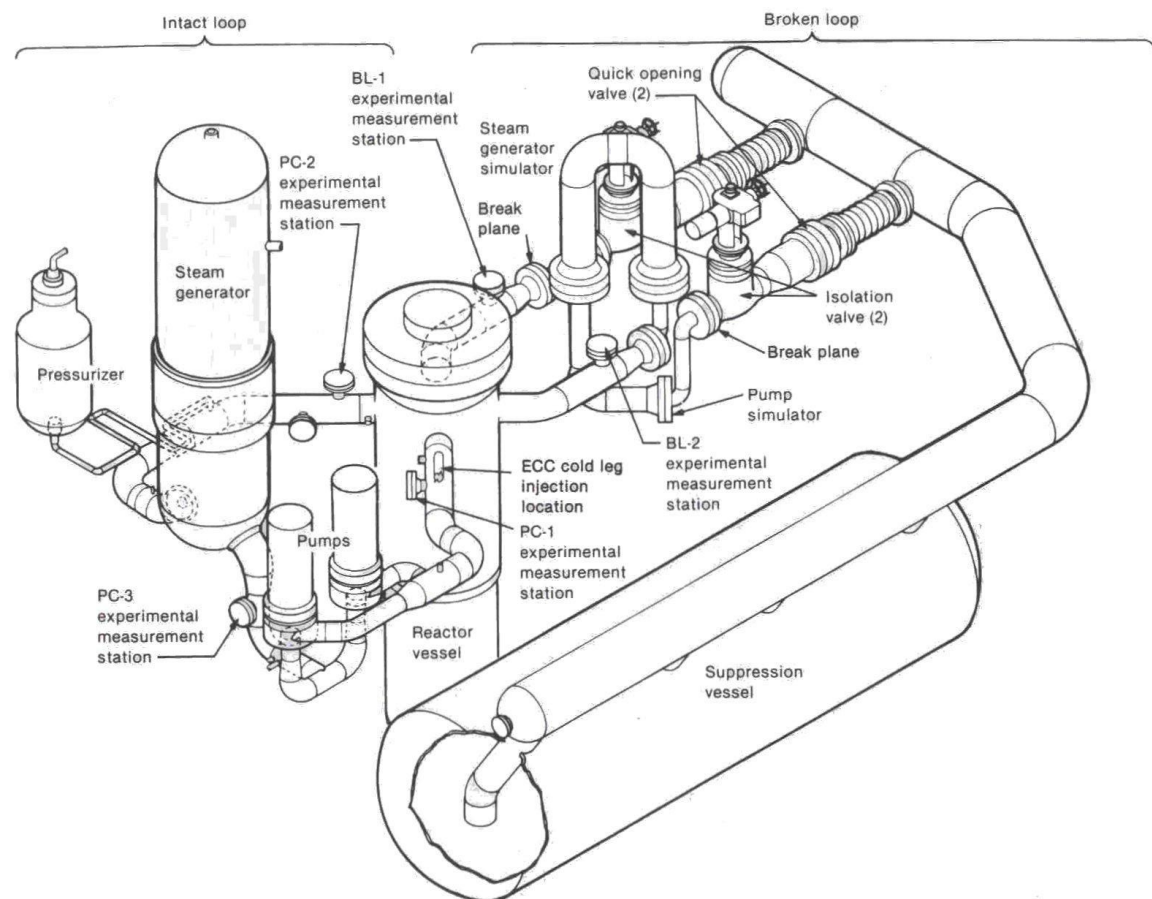
6.1. Obliczenia walidacyjne dla instalacji eksperymentalnej LOFT

Instalacja eksperymentalna LOFT jest jedną z niewielu instalacji doświadczalnych w których stosowano prawdziwe paliwo jądrowe zawierające dwutlenek uranu. Instalacja służyła do symulacji awarii zachodzących w reaktorach typu PWR. Instalacja składa się z następujących głównych składowych [42], [105]:

- Zbiornika reaktora, zawierającego rdzeń reaktora złożony z 1300 prętów paliwowych;
- Pętli obiegu pierwotnego bez rozerwania, na którą składają się rurociągi, wytwornica pary, stabilizator ciśnienia oraz dwie pompy cyrkulacyjne;

- Pętli obiegu pierwotnego z rozerwaniem, na którą składają się rurociągi, zwężki, dwa zawory odcinające oraz dwa zawory szybko otwierające się, które symulują gilotynowe rozerwanie;
- Systemu redukcji ciśnienia składającego się z rurociągów, zbiornika oraz spryskiwaczy;
- Systemu awaryjnego wtrysku chłodziwa na który składają się zbiornik akumulatora, wysokociśnieniowy system wtrysku awaryjnego (HPIS – „High Pressure Injection System”) oraz niskociśnieniowy system wtrysku awaryjnego (LPIS – „Low Pressure Injection System”).

Wszystkie składowe instalacji testowej zostały oprzyrządowane systemami pomiarowymi, aby zbierać jak największą liczbę danych pomiarowych. Schemat instalacji przedstawiono na rysunku 6.1.



Rys. 6.1. Schemat instalacji eksperymentalnej LOFT

Zbiornik reaktora w jak najdokładniejszy sposób odwzorowuje rzeczywiste zbiorniki stosowane w reaktorach typu PWR. Składa się ze szczeliny opadowej, dolnej komory

mieszania, rdzenia reaktora z paliwem UO_2 , górnej komory mieszania oraz elementów wewnątrz zbiornikowych takich jak dolna oraz górna płyta rdzenia. Wysokość elementów paliwowych wyniosła 1.676 metra, natomiast pastylki paliwowe posiadały długość 15.2 mm oraz średnicę 9.29 mm. Jedna pętla obiegu pierwotnego bez rozerwania w instalacji LOFT symuluje trzy pętle w reaktorze typu PWR i posiada rzeczywiste komponenty charakterystyczne dla reaktorów tego typu (wytwornica pary i stabilizator ciśnienia), natomiast pętla obiegu pierwotnego z rozerwaniem posiada odpowiednie komponenty pozwalające na symulowanie różnych typów awarii. Systemy bezpieczeństwa zostały zaprojektowane w sposób zapewniający możliwość podłączenia ich do różnych punktów w instalacji w zależności od badanego scenariusza awaryjnego. W teście L2-5, który został wybrany do analizy, systemy bezpieczeństwa zostały podłączone do zimnej nitki pętli bez rozerwania.

Model instalacji testowej LOFT został opracowany w kodzie TRACE stosując środowisko graficzne SNAP. Odwzorowane zostały wszystkie opisane powyżej składowe instalacji. Zbiornik reaktora został zamodelowany jako trójwymiarowy komponent zbiornika. Jest to jedyny komponent, dla którego można modelować warunki przepływu w trzech wymiarach. Pozwala to na o wiele dokładniejsze modelowanie procesów cieplno-przepływowych zachodzących w zbiorniku. Zbiornik zamodelowano stosując cztery składowe promieniowe – dwie dla rdzenia reaktora, jedna reprezentująca bocznik (ang. „reactor core bypass”) oraz jedna reprezentująca szczelinę opadową. W płaszczyźnie azymutalnej podzielono zbiornik na cztery części, zgodnie z liczbą nitek obiegu pierwotnego. W osiowej płaszczyźnie wydzielono 18 poziomów. Trzy dolne poziomy reprezentują dolną komorę mieszania, kolejne osiem reprezentuje wysokość rdzenia reaktora, następnie cztery komórki symulują przestrzeń między rdzeniem a króćcami nitek obiegu pierwotnego. Pozostałe komórki reprezentują górną komorę mieszania. Dedykowane komponenty zostały zastosowane również do modelowania stabilizatora ciśnienia, pomp cyrkulacyjnych, zbiorników akumulatora oraz systemu zrzutu ciśnienia. Pozostałe komponenty zamodelowano jako rurociągi, zawory oraz komponenty FILL i BREAK określające warunki brzegowe. Warunki wymiany ciepła określono poprzez zastosowanie struktur cieplnych i określania płaszczyzn wymiany ciepła (między komponentem a komponentem, bądź komponentem a warunkami środowiskowymi). Model składał się finalnie z 53 komponentów hydraulicznych oraz 36 struktur cieplnych. Dodatkowo opracowano również najważniejsze części systemu sterowania i zabezpieczeń, jak również sygnały wyzwajające działanie systemów bezpieczeństwa podczas awarii.

6.1.1. Kwalifikacja modelu instalacji LOFT – stan ustalony

Aby ocenić czy opracowany model w kompleksowy sposób odwzorowywał zarówno wymiary geometryczne jak i pozwalał na osiągnięcie warunków cieplno-przepływowych panujących w instalacji LOFT przeprowadzono kwalifikację modelu w warunkach stanu ustalonego. Pierwszym krokiem kwalifikacji jest porównanie rzeczywistych danych geometrycznych instalacji z objętościami komponentów modelujących składowe systemu w modelu przygotowanym w kodzie TRACE. Informacje na temat najważniejszych danych geometrycznych zestawiono w tabeli 6.1.

Tabela 6.1. Porównanie danych geometrycznych dla składowych instalacji LOFT z wartościami osiągniętymi w opracowanym modelu instalacji

	Parametr	Jednostka	Wartość zmierzona	Wartość w modelu TRACE	Błąd [%]
1	Objętość obiegu pierwotnego	m ³	7.69	7.61	1.0%
2	Objętość zbiornika reaktora	m ³	2.66	2.65	0.3%
3	Objętość w pętli bez rozerwania	m ³	3.67	3.64	0.8%
4	Objętość w pętli z rozerwaniem	m ³	1.36	1.31	3.5%
5	Objętość stabilizatora ciśnienia	m ³	0.963	0.956	0.7%
6	Objętość obiegu wtórnego	m ³	6.65	6.94	4.4%
7	Objętość zbiornika akumulatora	m ³	3.76	3.76	0.0%
8	Przekrój rurociągów w obiegu pierwotnym	m ²	0.0634	0.0634	0.0%
9	Średnia długość u-rurek w wytwornicy pary	m	5.17	5.14	0.6%
10	Wysokość paliwa jądrowego	m	1.676	1.676	0.0%

Dla ośmiu wielkości błąd nie przekroczył jednego procenta, co wskazuje na dobre odwzorowanie rzeczywistych wymiarów instalacji. Jedyne różnice występują w objętościach pętli z rozerwaniem oraz obiegu wtórnego, nie przekraczają one jednak pięciu procent. W przypadku obiegu pierwotnego (w tym pętli z zaworami symulującymi rozerwanie) należało mieć na względzie nie tylko objętości wszystkich komponentów, lecz również, zachowanie odpowiednich różnic wysokości w pętli, tak aby ostatecznie zimna i gorąca nitka została podłączona do zbiornika na odpowiedniej wysokości. Model obiegu wtórnego opracowany był na podstawie rysunków technicznych dostępnych w dokumentacji [105], stąd przy tak małych rozbieżnościach nie zdecydowano się na wprowadzanie zmian które nie byłyby zbieżne z danymi wynikającymi z dokumentacji technicznej. Ostatecznie uznano, więc przygotowany model za odwzorowujący wymiary i objętości instalacji doświadczalnej w wystarczający sposób.

Drugim krokiem kwalifikacji na etapie stanu ustalonego jest wykonanie obliczeń w celu sprawdzenia poprawności uzyskanych warunków cieplno-przepływowych w instalacji oraz

oceny ich stabilności. Porównanie dla trzynastu wybranych parametrów przedstawiono w tabeli 6.2. Wartości parametrów opisujących warunki cieplno-przepływowe powinny utrzymywać się na stabilnym poziomie w obliczeniach stanu ustalonego. Pozwala to na uzyskanie pewności, że uzyskane w trakcie symulacji stanów nieustalonych wyniki odwzorowują odpowiedzi modelowanych systemów, a nie wynikają z niestabilności numerycznych opracowanego modelu. Zgodnie z międzynarodową praktyką [42] przyjęto, że wartość parametru jest uznana za stabilną jeśli jej rozrzut jest mniejszy niż 1% w ciągu 100 sekund obliczeń stanu ustalonego. Największy zaobserwowany rozrzut w uzyskanych wynikach dla parametrów przedstawionych w tabeli 6.2 wynosił 0.2%.

Tabela 6.2. Kwalifikacja stanu ustalonego dla instalacji LOFT – porównanie wartości wybranych wielkości cieplno-przepływowych dla danych pomiarowych oraz uzyskanych w obliczeniach

Wielkość	Jednostka	Dane eksperymentalne	Zakres błędu pomiarowego	Wyniki obliczeń	Akceptowalny błąd	Rozbieżność
Moc	MW	36	1.2	36.0	2%	0%
Ciśnienie w gorącej nitce pętli bez rozerwania	MPa	14.94	0.1	14.93	0.1%	0%
Ciśnienie wyjściowe z wytwornicy		5.85	0.06	5.89		0%
Temperatura w gorącej nitce pętli bez rozerwania	K	589.5	4.2	590.9	0.5%	0%
Temperatura w zimnej nitce pętli bez rozerwania		556.4	4.3	557.2		0%
Temperatura w gorącej nitce pętli z rozerwaniem		561.9	4.3	557.2		0.1%
Temperatura w stabilizatorze ciśnienia		615.6	0.7	613.6		0.2%
Prędkość obrotowa pompy	rpm	1250	70	1453	1%	10.1%
Objętość wody w obiegu pierwotnym	kg	5330	300	5432	2%	0%
Wydatek przepływu w obiegu pierwotnym	kg/s	192.4	7.8	192.3	2%	0%
Wydatek przepływu wody zasilającej wytwornicę pary		19.1	0.4	18.8		0%
Poziom wody w stabilizatorze ciśnienia	m	1.14	0.03	1.17	4%	0%
Poziom wody w wytwornicy pary		2.9	0.1	2.94	3%	0%

Przedstawiona w ostatniej kolumnie rozbieżność reprezentuje ostateczny błąd po uwzględnieniu zakresu błędu pomiarowego. Spośród trzynastu parametrów dla dwunastu z nich - moc, wszystkie wartości temperatury, ciśnienia, wydatku przepływu, poziomu wody oraz

objętości wody w obiegu - uzyskany błąd wynosi maksymalnie 0.2% więc wszystkie te wartości mieszczą się w granicach akceptowalnego błędu. Rozbieżność dla prędkości obrotowej pompy wynosi 10%, jednak jednocześnie wydatek przepływu w obiegu pierwotnym który jest wartością ważniejszą z punktu widzenia warunków cieplno-przepływowych jest na tym samym poziomie co w eksperymencie. Na rozbieżność w prędkości obrotowej pompy mogą mieć wpływ zastosowane charakterystyki pracy pompy. Niemniej jednak zdecydowana większość parametrów jest zbieżna z wartościami eksperymentalnymi, obliczenia stanu ustalonego mogą więc zostać zakwalifikowane jako wykonane prawidłowo i odwzorowujące w wystarczającym stopniu warunki eksperymentalne.

6.1.2. Kwalifikacja modelu instalacji LOFT – stan nieustalony, obliczenia referencyjne

Analizowanym scenariuszem awaryjnym w instalacji LOFT był test o oznaczeniu L2-5, w którym analizowano gilotynowe rozerwanie rurociągu obiegu pierwotnego. W tym celu w instalacji zastosowano dwa zawory szybko otwierające się o średnicy równej średnicy rurociągów. W tabeli 6.3 przedstawiono sekwencję najważniejszych zdarzeń zachodzących w trakcie analizowanego scenariusza awaryjnego. Zestawiono czas zmierzony podczas przebiegu awarii w instalacji eksperymentalnej oraz czas uzyskany w obliczeniach referencyjnych wykonanych w kodzie TRACE.

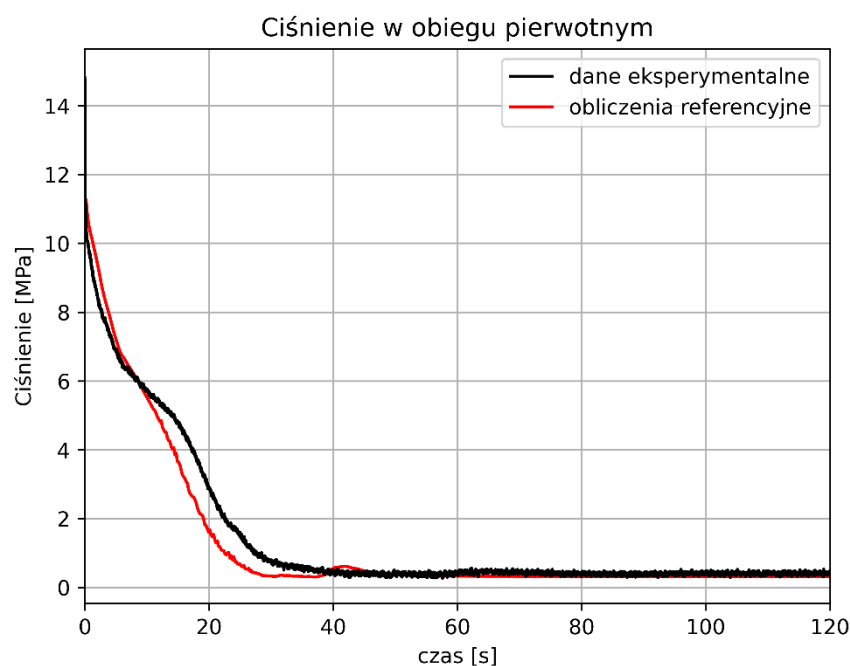
Tabela 6.3. Sekwencja zdarzeń w eksperymencie L2-5 w instalacji LOFT – porównanie danych eksperymentalnych oraz wyników obliczeń referencyjnych

Zdarzenie	Czas w eksperymencie (s)	Czas w obliczeniach (s)
Rozpoczęcie testu	0.0	0
Wyłączenie reaktora	0.24 ± 0.01	0.24
Pierwszy kryzys wrzenia	0.91 ± 0.01	1.62
Wyłączenie pomp obiegu pierwotnego	0.94 ± 0.01	0.94
Opróżnienie stabilizatora ciśnienia	15.4 ± 1.0	14.57
Rozpoczęcie wtrysku ze zbiornika akumulatora	16.8 ± 0.1	14.33
Rozpoczęcie zalewania rdzenia	22.7 ± 1.0	31.25
Rozpoczęcie wtrysku z HPIS	23.90 ± 0.02	23.98
Osiągnięcie maksymalnej temperatury koszulki	28.47 ± 0.02	8.67
Rozpoczęcie wtrysku z LPIS	37.32 ± 0.02	37.51
Akumulator opróżniony	49.6 ± 0.1	42.55
Rdzeń całkowicie zalany	65 ± 2.0	58.13
Zakończenie przepływu z LPIS	107.1 ± 0.4	120

Rozbieżności w uzyskanej sekwencji zdarzeń nie są znaczące. Kilkusekundowe rozbieżności między eksperymentem a obliczeniami są oczekiwane, nie jest możliwe perfekcyjne odwzorowanie warunków cieplno-przepływowych dla eksperymentów w dużej skali, w których badana jest duża liczba zjawisk. Jedynym zjawiskiem, dla którego zauważalna jest

duża rozbieżność jest czas osiągnięcia maksymalnej temperatury koszulki paliwowej. Wynika to z różnic w uzyskanych przebiegach temperatury koszulki. W eksperymencie obserwowane są dwa lokalne maksima – w dziesiątej oraz dwudziestej ósmej sekundzie przebiegu awarii. Większą wartość w eksperymencie posiada drugie maksimum, tymczasem w obliczeniach największą wartość temperatury koszulki obserwuje się w pierwszych dziesięciu sekundach, więc w zakresie czasowym pierwszego maksimum z eksperymentu. Ogólnie przebieg sekwencji zdarzeń wskazuje, że uruchomienie systemów bezpieczeństwa wyzwalane jest w odpowiednich momentach i w odpowiedniej kolejności, co dalej rzutuje na odpowiedni rozwój scenariusza awaryjnego.

Na rysunkach 6.2 – 6.5 przedstawiono przebiegi czterech istotnych parametrów wyjściowych, na podstawie których dokonana została ilościowa analiza uzyskanych wyników obliczeń referencyjnych. Zgodnie z opisem w poprzednich rozdziałach najistotniejsze będą masowe natężenie wypływu przez rozerwanie (rys. 6.3) oraz temperatura koszulki paliwowej (rys. 6.5). Niestety dokumentacja techniczna nie zawierała przebiegów zmian poziomu wody w rdzeniu. Parametr ten analizowany będzie w ramach walidacji, jednak bez wskazania przebiegu eksperymentalnego. Dwie dodatkowe wielkości wyjściowe dla których dokonano porównania danych pomiarowych oraz wyników obliczeń referencyjnych to ciśnienie w obiegu pierwotnym (rys. 6.2) oraz scałkowany wypływ z systemów bezpieczeństwa (rys. 6.4).



Rys. 6.2. Porównanie ciśnienia w obiegu pierwotnym dla danych eksperymentalnych w instalacji LOFT oraz obliczeń referencyjnych dla testu L2-5

Na rysunku 6.2 przedstawiono przebiegi ciśnienia w obiegu pierwotnym podczas eksperymentu oraz dla wyników obliczeń. Spadek ciśnienia w obiegu następuje szybciej w obliczeniach, jednak różnice nie są znaczące. Różnice te tłumaczą zauważalne w tabeli 6.3 niewielkie rozbieżności w czasach rozpoczęcia wtrysku z systemów bezpieczeństwa, których sygnały wyzwajające uzależnione są od wartości ciśnienia w obiegu. Kształt przebiegu ciśnienia jest jednak podobny do danych eksperymentalnych i mimo różnic w uzyskanych wartościach między 10 a 30 sekundą można uznać zgodność wyników obliczeń z danymi pomiarowymi za satysfakcjonującą.

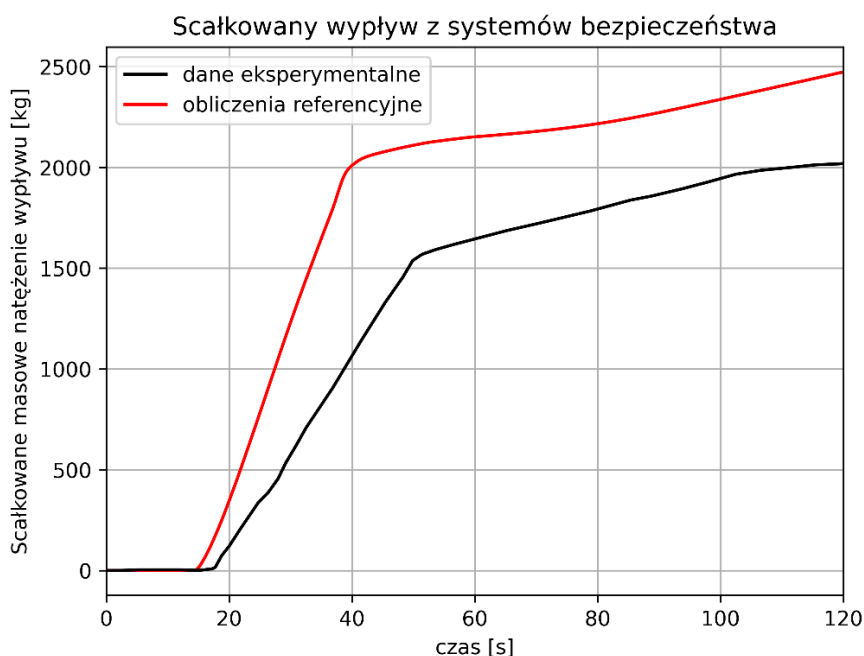


Rys. 6.3. Porównanie masowego natężenia przepływu przez rozerwanie dla danych eksperymentalnych w instalacji LOFT oraz obliczeń referencyjnych dla testu L2-5

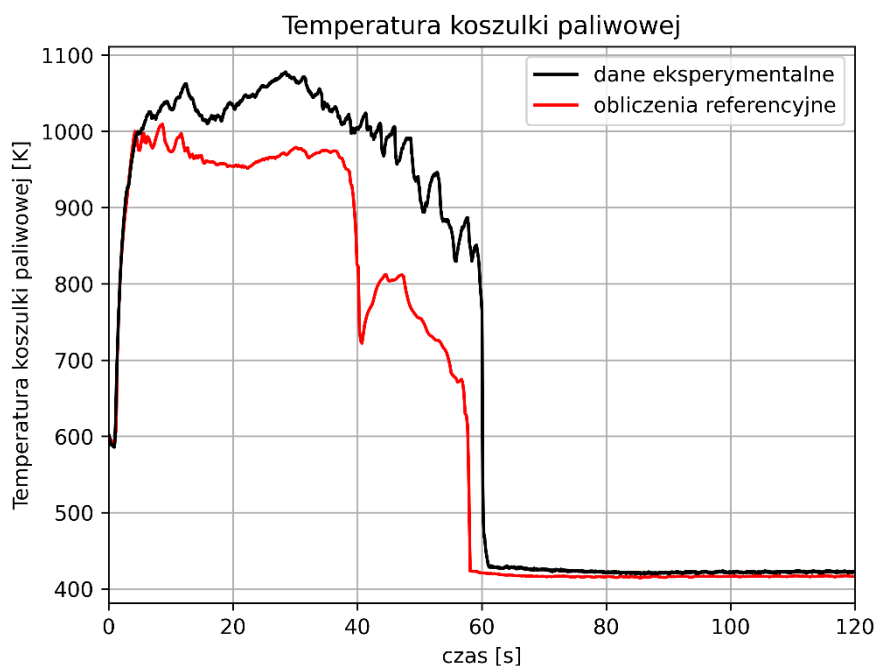
Na rysunku 6.3 przedstawiono przebiegi masowego natężenia przepływu przez rozerwanie podczas eksperymentu oraz dla wyników obliczeń. Zauważalne jest niewielkie przeszacowanie wielkości wypływu przez rozerwanie w obliczeniach wykonanych z zastosowaniem kodu TRACE przez pierwsze 40 sekund. Trend czasowy wydatku przepływu jest jednak zbliżony do danych eksperymentalnych.

Na rysunku 6.4 przedstawiono przebiegi skalkowanego sumarycznego przepływu z systemów bezpieczeństwa dla danych pomiarowych oraz wyników obliczeń. Zauważalny jest większy wypływ wody chłodzącej z systemów bezpieczeństwa (szczególnie z akumulatora) w obliczeniach, co następnie oddziałuje zarówno na wartości wydatku przepływu przez rozerwanie (rysunek 6.3) oraz na osiąganą temperaturę koszulki (rysunek 6.5). Objętość

zbiornika akumulatora jak zaznaczono w tabeli 6.1 była poprawnie zamodelowana, podobnie charakterystyki wypływu z systemów HPIS oraz LPIS, stąd rozbieżności wynikają prawdopodobnie z modelowania rurociągów systemów bezpieczeństwa, w których mogła znajdować się dodatkowa objętość wody.



Rys. 6.4. Porównanie skalkowanego wypływu z systemów bezpieczeństwa dla danych eksperymentalnych w instalacji LOFT oraz obliczeń referencyjnych dla testu L2-5



Rys. 6.5. Porównanie temperatury koszulki paliwowej dla danych eksperymentalnych w instalacji LOFT oraz obliczeń referencyjnych dla testu L2-5

Na rysunku 6.5 przedstawiono przebiegi temperatury koszulki paliwowej podczas eksperymentu oraz dla wyników obliczeń. Maksymalna temperatura koszulki paliwowej jest niedoszacowana w obliczeniach, jednak kształt przebiegu jest zbliżony do danych pomiarowych. Dodatkowo czas schłodzenia temperatury koszulki paliwowej jest przewidziany z dużą dokładnością. Jak wspomniano przy opisie rysunku 6.4 różnice w natężeniu wypływu z systemów bezpieczeństwa wpływają bezpośrednio na ilość wody dostępną do chłodzenia rdzenia, stąd następuje szybsze zalanie rdzenia i temperatura koszulki paliwowej już w czterdziestej sekundzie zaczyna bardziej gwałtownie opadać.

Tabela 6.4. Kwalifikacja wyników obliczeń referencyjnych

Wielkość	Jednostka	Dane eksperymentalne	Obliczenia TRACE	Błąd
Maksymalne natężenie przepływu	kg/s	777.0	860.9	11%
Scałkowany wypływ przez rozerwanie w momencie rozpoczęcia wtrysku z akumulatora	kg	3768.0	4491.4	19%
Scałkowany wypływ przez rozerwanie po upływie 100 s	kg	5160.0	6376.8	24%
Stosunek ciśnienia w stabilizatorze do ciśnienia w obiegu po 10 s	-	2.01	1.86	-8%
Czas, w którym ciśnienie w stabilizatorze zrównuje się z ciśnieniem w obiegu	s	38.0	28.84	-24%
Czas wystąpienia maksymalnej temperatury koszulki	s	28.47	8.67	-70%
Maksymalna temperatura koszulki	K	1078.0	1009.5	-6%
Czas całkowitego zalania rdzenia	s	65.0	58.13	-11%

W tabeli 6.4 przedstawiono porównanie wyników obliczeń oraz danych eksperymentalnych dla ośmiu skalarnych wielkości opisujących natężenie wypływu (w tym scałkowane) przez rozerwanie, ciśnienie w stabilizatorze i obiegu oraz przebiegu zmian temperatury koszulki paliwowej w czasie i przebiegu propagacji frontu zalewania. Najmniejsze rozbieżności uzyskano dla wartości maksymalnej temperatury koszulki oraz stosunku ciśnień w stabilizatorze i obiegu pierwotnym, gdzie rozbieżność nie przekraczała 10%. Największa rozbieżność dotyczyła czasu wystąpienia maksymalnej temperatury koszulki paliwowej, co zostało omówione powyżej w dyskusji dotyczącej sekwencji zdarzeń. Dla trzech parametrów określających scałkowany wypływ przez rozerwanie oraz czas, w którym ciśnienie w stabilizatorze zrównuje się z ciśnieniem w obiegu rozbieżność wynosiła od 19 do 24%. Rozbieżności są więc zauważalne, jednak przy tak dynamicznie przebiegającym scenariuszu jak LBLOCA są one do zaakceptowania. Należy również mieć na uwadze, że obliczenia referencyjne są jedynie krokiem do celu, czyli do przeprowadzenia obliczeń

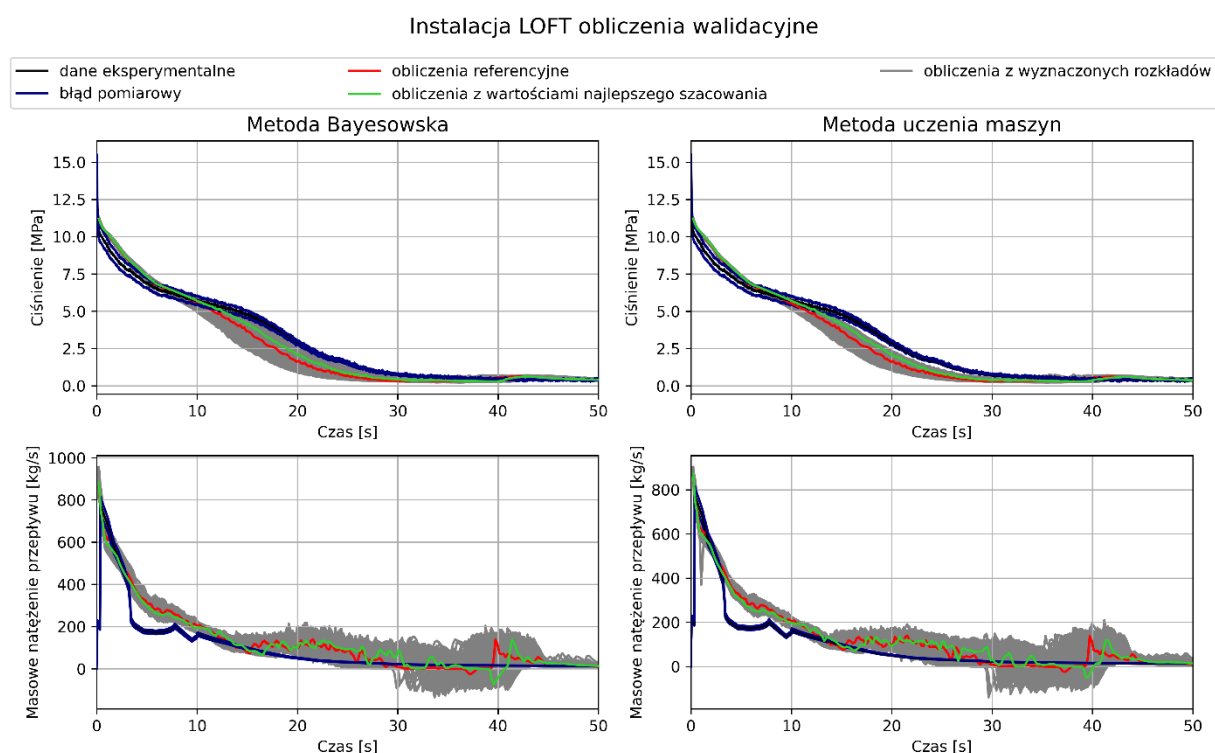
walidacyjnych, pozwalających na ocenę poprawności wyznaczonych w poprzednim rozdziale rozkładów gęstości prawdopodobieństwa wybranych parametrów. Wymagane jest więc uzyskanie wyników obliczeń referencyjnych na satysfakcjonującym poziomie, które nie będą znacząco odbiegać od danych eksperymentalnych, gdyż jak zaznaczano w rozdziale drugim, obliczenia referencyjną są istotną bazą do obliczeń niepewności. Uzyskane wyniki można zarówno ilościowo jak i jakościowo ocenić jako satysfakcjonujące i pozwalające na wykonanie obliczeń walidacyjnych stosując metodę propagacji niepewności wejściowych parametrów modelowych kodu.

6.1.3. Walidacja rozkładów gęstości prawdopodobieństwa parametrów – obliczenia dla instalacji LOFT (element piąty metodyki)

Zgodnie z opisem w [rozdziale 3.5](#) walidacji wyznaczonych w elemencie czwartym rozkładów gęstości prawdopodobieństwa parametrów wejściowych opisujących modele stosowane w kodzie obliczeniowym dokonuje się wykonując obliczenia dla instalacji eksperymentalnej w której badane są jednocześnie wszystkie zjawiska, charakteryzowane przez te parametry wejściowe. Zgodnie z krokami metodyki przed przeprowadzeniem obliczeń walidacyjnych należy ustalić inne źródła niepewności, a następnie wykonać obliczenia pozwalające na oszacowanie niepewności wyjściowych stosując metodę propagacji niepewności wejściowych. Zdecydowano jednak, że na tym etapie, przy pierwszym wdrożeniu metodyki, pominięte zostaną inne źródła niepewności, dopiero w obliczeniach dla reaktora energetycznego (rozdział 6.2) zostaną one wzięte pod uwagę. Powodowane jest to chęcią bezpośredniej obserwacji wpływu wyznaczonych rozkładów gęstości prawdopodobieństwa dziewięciu badanych parametrów, bez wprowadzania dodatkowych parametrów które mogłyby mieć większy wpływ na niepewności wyjściowe.

Na podstawie wyników kwantyfikacji niepewności parametrów wejściowych przedstawionych w rozdziale 5 uzyskano dwa zestawy rozkładów dziewięciu parametrów – uzyskane stosując metodę Bayesowską oraz metodę uczenia maszyn. Z każdego zestawu losowano 1000 razy wartości parametrów stosując wyznaczone w rozdziale piątym rozkłady gęstości prawdopodobieństwa. Wartości te następnie wprowadzano do pliku wsadowego kodu TRACE opisującego model instalacji LOFT (opisany i kwalifikowany powyżej). Pliki wsadowe do obliczeń referencyjnych oraz obliczeń niepewności różnić się będą, więc jedynie wartościami parametrów CHM12, CHM22, C1RC2, C2RC2, 1009, 1011, 1014, 1031 oraz 1033 które omówiono w rozdziale piątym.

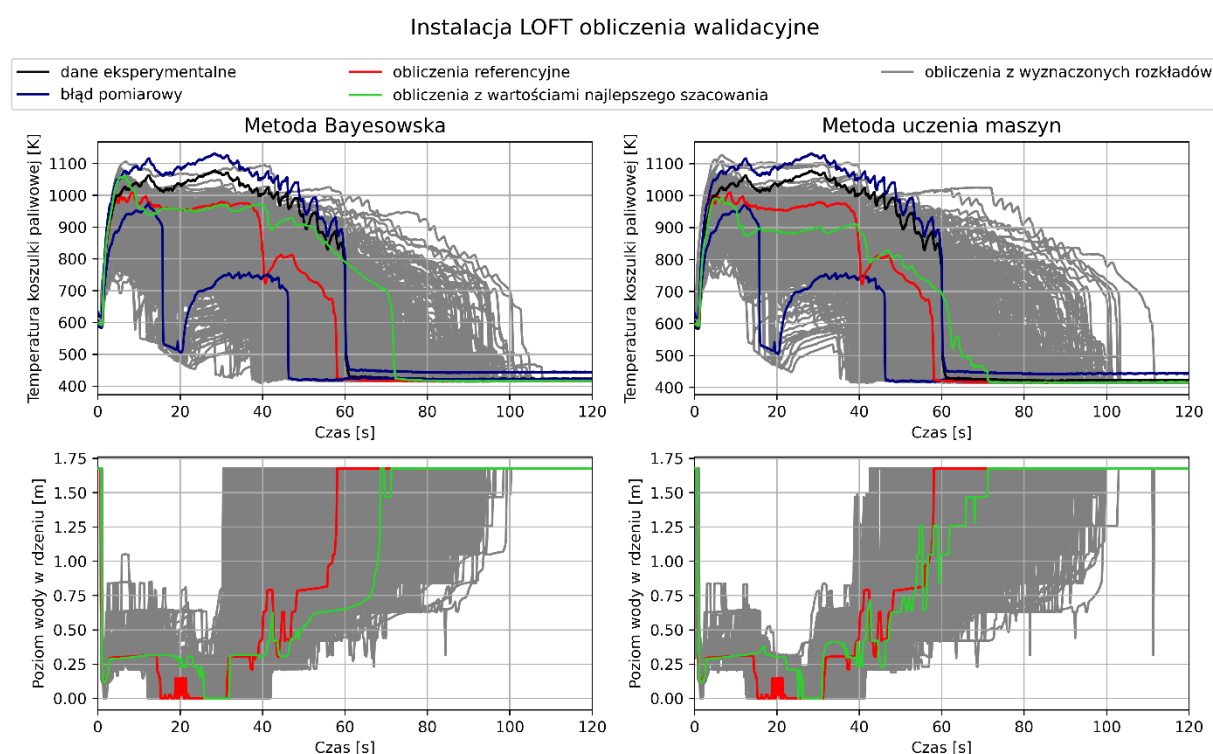
Po przygotowaniu 1000 plików wsadowych dla każdego z zestawu parametrów, wykonano łącznie 2000 obliczeń (po 1000 dla każdej z metod na podstawie której określono rozkłady parametrów). Obliczenia dla jednego przypadku trwały około 30 minut, stąd przy wykorzystaniu sześciu rdzeni komputera całkowity czas obliczeń wyniósł osiem dni. Dodatkowo dla każdej z metod wykonano obliczenia z wartościami najlepszego szacowania parametrów, określonymi w rozdziale piątym.



Rys. 6.6. Zakresy niepewności ciśnienia w obiegu oraz masowego natężenia przepływu w obliczeniach walidacyjnych dla instalacji LOFT stosując rozkłady parametrów wejściowych uzyskane z zastosowaniem dwóch opracowanych metod

Na rysunku 6.6 przedstawiono wyniki obliczeń walidacyjnych dla ciśnienia w obiegu pierwotnym oraz dla masowego natężenia przepływu. Przebiegi przedstawiono do pięćdziesiątej sekundy, ponieważ, jak można zaobserwować na rysunkach 6.2 oraz 6.3, dalej wartości równe są ciśnieniu atmosferycznemu oraz zeru w przypadku masowego natężenia przepływu. Dla przebiegu ciśnienia obliczenia z wartościami najlepszego szacowania dla obu metod pozwalają na uzyskanie rezultatów nieznacznie bliższych danym eksperymentalnym w porównaniu do obliczeń referencyjnych. Zakres niepewności dla metody Bayesowskiej jest szerszy i pozwala na objęcie danych eksperymentalnych. Zakres niepewności dla metody uczenia maszyn jest węższy i nie pokrywa danych eksperymentalnych, jednakże rozbieżności nie są duże. W przypadku masowego natężenia przepływu obliczenia najlepszego szacowania

są w obu przypadkach jedynie nieznacznie lepiej dopasowane do danych eksperymentalnych niż obliczenia referencyjne. Dodatkowo w obu przypadkach poza przedziałem od piątej do dziesiątej sekundy zakresy niepewności w pełni pokrywają dane eksperymentalne. Uzyskane zakresy niepewności są najszersze między 30 a 40 sekundą, gdy ustala się przepływ dwufazowy. Ogólnie, jednakże uzyskane zakresy niepewności można uznać za poprawnie wyznaczone, stąd można wnioskować, że zaproponowane rozkłady gęstości parametrów CHM12, CHM22, C1RC2 oraz C2RC2 pozwalają na uzyskanie wyników pokrywających dane eksperymentalne.



Rys. 6.7. Zakresy niepewności temperatury koszulki paliwowej oraz poziomu wody w rdzeniu w obliczeniach walidacyjnych dla instalacji LOFT stosując rozkłady parametrów wejściowych uzyskane z zastosowaniem dwóch opracowanych metod

Na rysunku 6.7 przedstawiono wyniki obliczeń walidacyjnych dla temperatury koszulki paliwowej oraz poziomu wody w rdzeniu. Dla przebiegu zmian temperatury koszulki paliwowej uzyskane zakresy są bardzo szerokie, pozwalając na pokrycie danych eksperymentalnych. Maksymalna temperatura koszulki nie jest przeszacowana zarówno w obliczeniach stosujących metodę Bayesowską jak i metodę uczenia maszyn. Maksymalne wartości temperatury koszulki paliwowej nie przekraczają ani kryterium akceptacji (1473 K) ani górnego zakresu błędu pomiarowego. Dolny dopuszczalny zakres temperatury koszulki paliwowej jest dokładniej oszacowany przez obliczenia wykonane stosując metodę uczenia

maszyn. Czas, w którym temperatura osiąga wartość minimalną zawiera się w szerokim zakresie od 40 do 100 sekund w obu przypadkach. Dodatkowo można zaobserwować pojedyncze wartości odstające w metodzie uczenia maszyn. Uzyskiwane zakresy niepewności są węższe dla obliczeń w których wykorzystywano rozkłady gęstości prawdopodobieństwa parametrów uzyskanych stosując metodę uczenia maszyn. Obliczenia z wartościami najlepszego szacowania parametrów wejściowych modelu są bardziej zgodne z danymi eksperymentalnymi przy zastosowaniu metody Bayesowskiej. Mimo, że przeszacowany został moment osiągnięcia minimalnej temperatury to kształt przebiegu zmian temperatury jest prawidłowy. W przypadku obliczeń dla metody uczenia maszyn, przebieg obliczeń dla wartości najlepszego szacowania bardziej przypomina obliczenia referencyjne niż dane eksperymentalne.

Rozkłady niepewności dla przebiegu poziomu wody w rdzeniu wykazują duże rozbieżności w odniesieniu do wyznaczania czasu całkowitego zalania rdzenia. Węższy zakres zaobserwować można w obliczeniach wykonanych z zastosowaniem metody uczenia maszyn. W obu przypadkach można zaobserwować, że przebieg poziomu wody w rdzeniu podąża za zmianami temperatury koszulki paliwowej, gdyż czas całkowitego zalania rdzenia jest skorelowany z czasem, gdy temperatura koszulki paliwowej osiąga wartość minimalną.

Oceniając uzyskiwane zakresy niepewności przyjmuje się, że nie powinny one być nadmiernie szerokie oraz powinny pozwalać na pokrycie danych eksperymentalnych. W przypadku parametrów CHM12, CHM22, C1RC2 oraz C2RC2 uzyskane zakresy niepewności dla masowego natężenie przepływu spełniają te kryteria dla obu metod, pomimo że przez krótki okres czasu (między 5 a 10 sekundą) dane eksperymentalne nie są pokryte. Rozbieżności w tym krótkim przedziale czasowym nie są jednak duże. Rozkłady uzyskane stosując obie opracowane metody kwantyfikacji niepewności są więc poprawne i mogą zostać zastosowane do dalszych obliczeń niepewności dla elektrowni jądrowej typu PWR. W przypadku pięciu parametrów które opisywały wymianę ciepła w rdzeniu ich wpływ na uzyskiwane wyniki obserwowalny był dla temperatury koszulki paliwowej oraz poziomu wody w rdzeniu. Uzyskane zakresy niepewności są szerokie jednak pozwalają na pokrycie danych eksperymentalnych. Ogólnie dla obu badanych wielkości wyjściowych węższe zakresy uzyskuje się dla parametrów uzyskanych kwantyfikując niepewności metodą uczenia maszyn. Stąd te rozkłady gęstości prawdopodobieństwa pozwalają na uzyskanie lepszych rezultatów. Należy zwrócić jednak uwagę, że rozbieżności między metodami nie były duże, więc w dalszych ostatecznych obliczeniach dla reaktora energetycznego wciąż stosowane będą zestawy parametrów

uzyskiwane z zastosowaniem obu metod. Zestaw wartości najlepszego szacowania uzyskany stosując metodę Bayesowską powinien być stosowany w przyszłych badaniach.

6.2. Obliczenia niepewności dla modelu elektrowni jądrowej typu PWR

Finalna walidacja rozkładów gęstości prawdopodobieństwa przeprowadzana jest poprzez wykonanie obliczeń niepewności dla modelu elektrowni jądrowej, zgodnie z krokiem szóstym metodyki. Obliczenia wykonywane są dla scenariusza określonego w rozdziale czwartym, czyli dużego rozerwania rurociągu obiegu pierwotnego w reaktorze typu PWR. Model elektrowni jądrowej opracowano na podstawie pliku wsadowego dla trzy pętlowej elektrowni typu PWR produkcji Westinghouse opisanego w instrukcji kodu TRAC-M [106]. TRAC-M jest poprzednikiem wykorzystywanego w obecnej pracy do wykonywania obliczeń kodu TRACE. Instrukcja użytkowania kodu TRAC-M zawiera opis pliku wsadowego dla modelu trzypętlowej elektrowni jądrowej, co wraz z przypisami stanowiło przykład zastosowania kodu i wykorzystania dobrych praktyk. Uwzględniając zmiany, które należało wprowadzić, aby model był zgodny z funkcjonalnościami kodu TRACE, oraz dobre praktyki modelowania w kodzie TRACE, opracowano nowy model tej elektrowni w kodzie TRACE. Wprowadzone zmiany dotyczyły między innymi struktur cieplnych, modelowania zbiornika reaktora, modelowania wytwornicy pary, systemów bezpieczeństwa oraz systemu sterowania i zabezpieczeń. Bardziej szczegółowy opis wprowadzonych zmian oraz całego modelu znajduje się w artykułach [107], [108].

Opis modelowanej elektrowni postanowiono połączyć z opisem zastosowanych rozwiązań modelowych w kodzie TRACE, aby w klarowny sposób przedstawić cały model. Najważniejszymi komponentami modelowanej elektrowni jądrowej są:

- Zbiornik reaktora jądrowego – modelowany przez trójwymiarowy komponent zbiornika, podzielony na dwie składowe promieniowe, sześć azymutalnych oraz dwadzieścia trzy składowe osiowe. Dwie składowe promieniowe reprezentują szczelinę opadową oraz rdzeń reaktora. Części azymutalne odpowiadają liczbie nitek (zimna i ciepła) w trzech pętlach reaktora. W płaszczyźnie osiowej pięć komórek reprezentuje dolną komorę mieszania, dwanaście rdzeń reaktora, a pozostałe górną część zbiornika oraz górną komorę mieszania. Sześć struktur cieplnych (po jednej na każdą część azymutalną) reprezentuje paliwo jądrowe w rdzeniu reaktora. Dziewięć struktur cieplnych reprezentuje elementy i płyty wewnątrz-zbiornikowe. Kolejne trzydzieści struktur cieplnych modelowało kosz rdzenia reaktora oraz ścianki zbiornika reaktora.

- Trzy pętle obiegu pierwotnego modelowane były oddzielne. Każda pętla składa się z rurociągów gorącej nitki, zimnej nitki, u-rurek wytwornicy pary oraz pompy cyrkulacyjnej obiegu pierwotnego. Dla pomp cyrkulacyjnych zastosowano modele pomp z wbudowanymi charakterystykami dla reaktorów typu Westinghouse. W pierwszej pętli obiegu pierwotnego zamodelowano gilotynowe rozerwanie rurociągu poprzez dodanie dwóch zaworów o średnicy odpowiadającej średnicy rurociągu oraz dwóch komponentów BREAK które reprezentują warunki panujące w obudowie bezpieczeństwa reaktora. W trzeciej pętli obiegu pierwotnego zamodelowano stabilizator ciśnienia, wraz z jego rurociągami oraz zaworami zrzutowymi. Stabilizator modelowany jest poprzez rurociąg składający się z 13 komórek oraz komponent PRIZER które ustala zakładane ciśnienie w stabilizatorze ciśnienia w trakcie obliczeń stanu ustalonego.
- Obieg wtórny składa się z trzech wytwornic pary, rurociągów wody zasilającej oraz rurociągów parowych wraz z zaworami upustowymi i zrzutowymi. Komponenty te modelowane są poprzez rurociągi, zawory oraz komponenty BREAK i FILL. Istotnym elementem obiegu wtórnego wytwornicy pary jest separator pary, dla którego w kodzie TRACE dedykowany jest specjalny komponent, który działa jako separator oraz osuszacz pary.
- System regulacji chemicznej i objętości chłodziwa reaktora modelowany w uproszczeniu przez trzy komponenty FILL pozwalające na ustalenie dodatkowego przepływu chłodziwa w zależności od ciśnienia.
- System awaryjnego zalewania rdzenia reaktora składający się ze zbiornika akumulatora, wysokociśnieniowego systemu wtrysku awaryjnego (HPIS – „High Pressure Injection System”) oraz niskociśnieniowego systemu wtrysku awaryjnego (LPIS – „Low Pressure Injection System”).
- System zabezpieczeń i kontroli wraz z sygnałami wyzwalającymi, co pozwala na sterowanie ciśnieniem i poziomem wody w stabilizatorze ciśnienia, działaniem grzałek w stabilizatorze ciśnienia, poziomem wody w wytwornicy pary, ustalaniem warunków przepływu wody zasilającej oraz działaniem wszystkich zaimplementowanych w modelu zaworów.

6.2.1. Kwalifikacja modelu elektrowni jądrowej – stan ustalony

Kwalifikacja modelu elektrowni jądrowej przeprowadzona została jedynie dla wyników obliczeń stanu ustalonego. Dane geometryczne są zbieżne z danymi przedstawionymi w źródłowym pliku wsadowym, stąd brak innych danych do porównania. W dokumentacji przedstawiono jednak oczekiwane wartości najważniejszych parametrów fizycznych reprezentujących stan elektrowni jądrowej, co pozwoliło na kwalifikację obliczeń stanu ustalonego dla modelu opracowanego w kodzie TRACE. Zestawienie danych źródłowych z dokumentacji oraz uzyskanych obliczeniowo przedstawiono w tabeli 6.5. Obliczenia stanu ustalonego wykonano dwuetapowo. W pierwszym etapie zastosowano wewnętrzne możliwości kodu TRACE pozwalające na użycie regulatorów proporcjonalno-całkująco-różniczkujących, aby ustalić między innymi prędkości pomp cyrkulacyjnych. Dodatkowo opracowane systemy kontrolne pozwoliły na ustalenie zakładanych poziomów wody w stabilizatorze ciśnienia i wytwornicy pary oraz przepływu pary w głównym rurociągu parowym. Tak ustalone warunki były podstawą do wykonania obliczeń trwających 200 sekund w czasie których nie działają dodatkowe systemy kontrolne ani regulatory, aby sprawdzić stabilność uzyskanych wyników. Uzyskane wyniki obliczeniowe były stabilne, różnice między największą a najmniejszą wartością parametrów określonych w tabeli 6.5 nie przekraczały 0.5%.

Tabela 6.5. Kwalifikacja stanu ustalonego dla modelu elektrowni jądrowej – porównanie wartości wybranych wielkości cieplno-przepływowych dla danych źródłowych oraz uzyskanych w obliczeniach

Wielkość	Jednostka	Dane źródłowe	Wyniki obliczeń	Akceptowalny błąd	Rozbieżność
Moc	MW	2300	2301	2 %	0.1 %
Ciśnienie w stabilizatorze ciśnienia	MPa	15.51	15.53	0.1 MPa	0.02 MPa
Ciśnienie wyjściowe z wytwornicy		5.52	5.55		0.03 MPa
Temperatura w gorącej nitce	K	591.1	589.3	0.5 %	0.3 %
Temperatura w zimnej nitce		559.1	556.6		0.4 %
Temperatura wody zasilającej		500.6	500.6		0 %
Temperatura pary w rurociągu parowym		542.1	543.7		0.3 %
Wydatek przepływu w obiegu pierwotnym	kg/s	4258.7	4257.7	2 %	0%
Wydatek przepływu wody zasilającej wytwornicę pary		424.6	423.7		0.2 %

Dla wszystkich badanych parametrów uzyskana rozbieżność między wynikami obliczeń a danymi źródłowymi nie przekracza akceptowalnego błędu. Największe rozbieżności zaobserwowano dla warunków w stabilizatorze ciśnienia oraz zimnej nitce pętli obiegu

pierwotnego, jednak wciąż są one na akceptowalnym poziomie. Obliczenia na etapie stanu ustalonego można zaklasyfikować jako stabilne, wykonane prawidłowo i odwzorowujące w wystarczającym stopniu dane źródłowe.

6.2.2. Kwalifikacja modelu instalacji elektrowni jądrowej – stan nieustalony, obliczenia referencyjne

Awarią analizowaną w ramach obliczeń walidacyjnych jest gilotynowe rozerwanie rurociągu zimnej nitki obiegu pierwotnego. Podstawowe założenia dla rozpatrywanej awarii to między innymi utrata zasilania zewnętrznego, co jest zgodne z wymaganiami rozporządzenia o analizach [6]. Jako kryterium pojedynczego uszkodzenia (wymagane dla awarii projektowych przez rozporządzenie o analizach oraz praktykę światową) przyjęto ograniczenie przepływu z wysokociśnieniowego i niskociśnieniowego systemu wtrysku awaryjnego. Podstawą projektową jest praca trzech pomp dla każdego z tych systemów. W założeniach do awarii przyjęto, że jedynie jedna pompa jest dostępna, a dwie są niesprawne. Poza niezadziałaniem jednej z pomp w ramach kryterium pojedynczego uszkodzenia, zakładano niedostępność jednej z pomp w ramach przeprowadzania prac konserwacyjnych, zgodnych ze specyfikacją techniczną. Dodatkowo założono, że ograniczony jest przepływ z pomocniczego systemu wody zasilającej. Zmiany w modelu obejmowały dostosowanie czasu wyzwolenia niektórych sygnałów oraz określenie zmian ciśnienia w obudowie bezpieczeństwa w trakcie awarii.

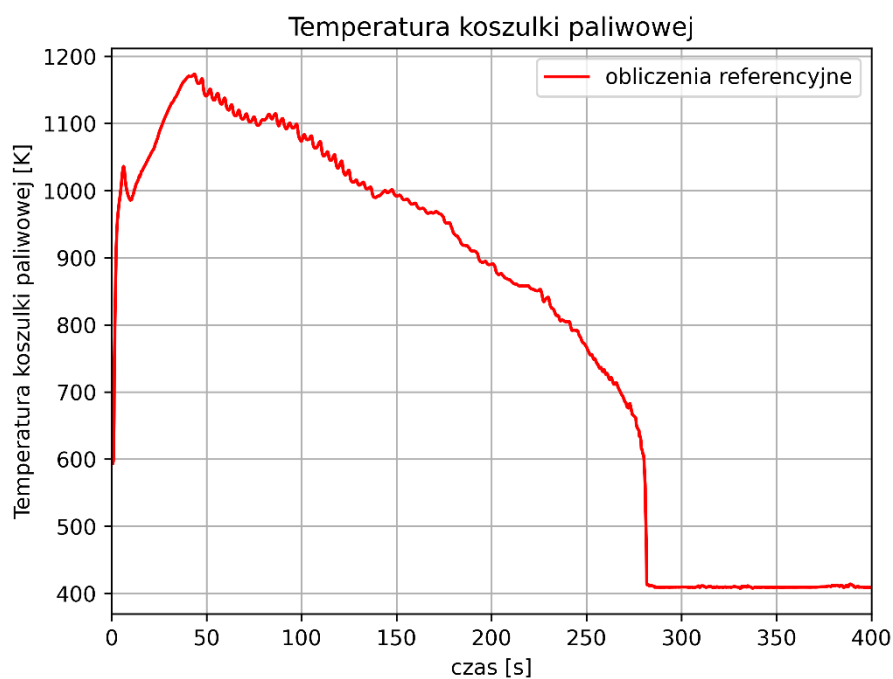
Tabela 6.6. Sekwencja zdarzeń dla awarii typu LBLOCA w reaktorze typu PWR

Zdarzenie	Czas w danych źródłowych [s]	Czas w obliczeniach [s]
Rozpoczęcie testu	0	0
Wyłączenie reaktora	0.05	0.9
Wyłączenie pomp obiegu pierwotnego	0.1	0.9
Wyzwolenie sygnału SI (Safety Injection - wtrysk z systemów bezpieczeństwa)	0.73	3.92
Rozpoczęcie wtrysku ze zbiornika akumulatora	2.7 (Pętla z rozerwaniem)	2.51
	11.5 (Pętla bez rozerwania)	8.13
Rozpoczęcie wtrysku z HPIS	31.2	32.11
Akumulator opróżniony	36.9 (Pętla z rozerwaniem)	55.18
	44.1 (Pętla bez rozerwania)	64.3
Rozpoczęcie zalewania rdzenia	39.3	35.6
Rozpoczęcie wtrysku z LPIS	43.2 (Pętla z rozerwaniem)	48.98
	44.1 (Pętla bez rozerwania)	48.98
Osiągnięcie maksymalnej temperatury koszulki	123.1	43.63

Dane źródłowe w dokumentacji kodu TRAC-M nie zawierały opisu scenariuszy awaryjnych. W artykule [107] zawarto odniesienia do raportu bezpieczeństwa dla elektrowni która jest modelowana. Dane z raportu bezpieczeństwa były jednak używane na pozwoleniu, które

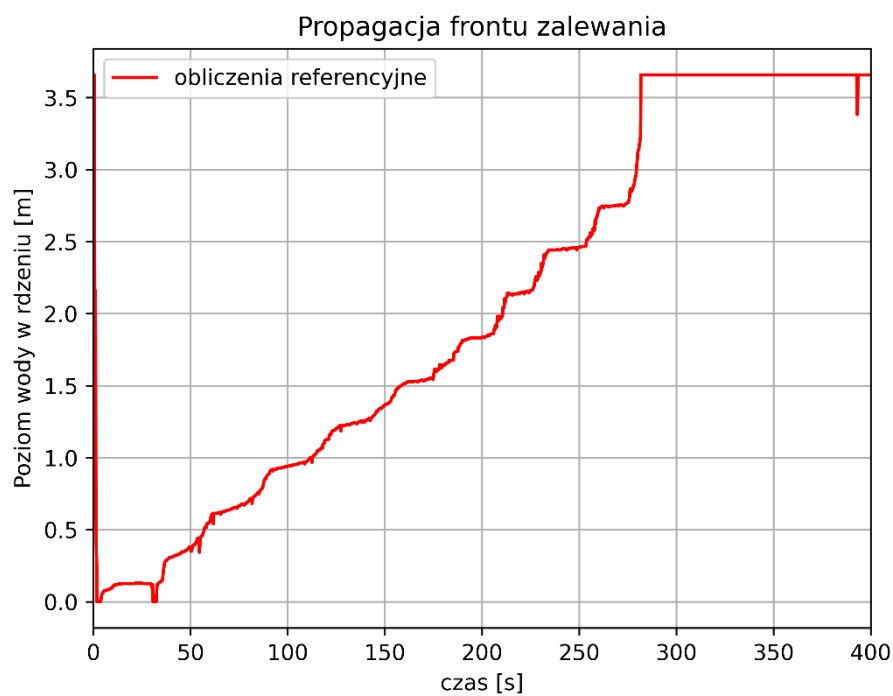
uzyskało Brookhaven National Laboratory, którego pracownicy byli współautorami publikacji. Ze względu na potencjalne naruszenie warunków stosowania tych danych nie zostaną one wykorzystane w poniższej rozprawie. Przedstawione zostaną jedynie informacje dotyczące sekwencji zdarzeń dla danej awarii, które dostępne były w ramach publicznie dostępnych raportów technicznych. Należy jednak pamiętać, że sekwencja zdarzeń przygotowana była dla obliczeń wykonywanych stosując założenia zachowawcze, natomiast obliczenia wykonywane w ramach rozprawy opierają się na podejściu najlepszego szacowania wraz z oceną niepewności. Porównanie danych źródłowych dotyczących przebiegu awarii oraz uzyskanych w obliczeniach stosując kod TRACE przedstawiono w tabeli 6.6.

Uzyskane rozbieżności między sekwencją zdarzeń w danych źródłowych oraz w wykonanych obliczeniach nie są znaczące. Wyzwolenie sygnału SI (sygnalizacja możliwości rozpoczęcia wtrysku z systemów bezpieczeństwa HPIS oraz LPIS) następuje trzy sekundy później w wykonanych obliczeniach, co nieznacznie wpływa na opóźnienie w rozpoczęciu wtrysku z systemów HPIS oraz LPIS. Rozpoczęcie wtrysku ze zbiorników akumulatorów następuje nieznacznie wcześniej w obliczeniach co oznacza, że spadek ciśnienia w obiegu pierwotnym następował szybciej niż w danych źródłowych. Wtrysk z akumulatorów trwał również około 20 sekund dłużej niż w obliczeniach źródłowych. Ogólnie przebieg sekwencji zdarzeń wskazuje, że uruchomienie systemów bezpieczeństwa wyzwalane jest w odpowiednich momentach i w odpowiedniej kolejności, co dalej rzutuje na odpowiedni rozwój scenariusza awaryjnego. Największą rozbieżność można zaobserwować dla czasu osiągnięcia maksymalnej temperatury koszulki paliwowej, gdzie w obliczeniach wartość ta wystąpiła zdecydowanie wcześniej. Podobne zachowanie obserwowano w obliczeniach dla instalacji testowej LOFT, wówczas jednak obserwowano dwa lokalne maksima, z których wartość wcześniejsza była większa, tymczasem w obliczeniach dla elektrowni maksimum jest jedno, a następnie temperatura koszulki powoli opada do osiągnięcia wartości temperatury minimalnej po 480 sekundach. Przebieg temperatury koszulki paliwowej dla najgorętszego pręta w rdzeniu zaprezentowano na rysunku 6.8. Uzyskana wartość maksymalnej temperatury koszulki równa 1173.5 K jest o 300 K mniejsza od limitu bezpieczeństwa wynoszącego 1473 K. Margines bezpieczeństwa jest więc znaczący co jest zgodne z celami i założeniami stosowania podejścia najlepszego szacowania wraz z oceną niepewności opisanego w [rozdziale 1.1](#). Temperatura koszulki paliwowej osiąga wartość minimalną po 283 sekundach.



Rys. 6.8. Przebieg temperatury koszulki paliwowej dla obliczeń referencyjnych dla modelu elektrowni jądrowej

Na rysunku 6.9 przedstawiono przebieg poziomu wody w rdzeniu. Rdzeń jest odkryty po dwóch sekundach, a faza napełniania trwa do 30 sekundy. Zalewanie rdzenia rozpoczyna się w 32 sekundzie i trwa 250 sekund.



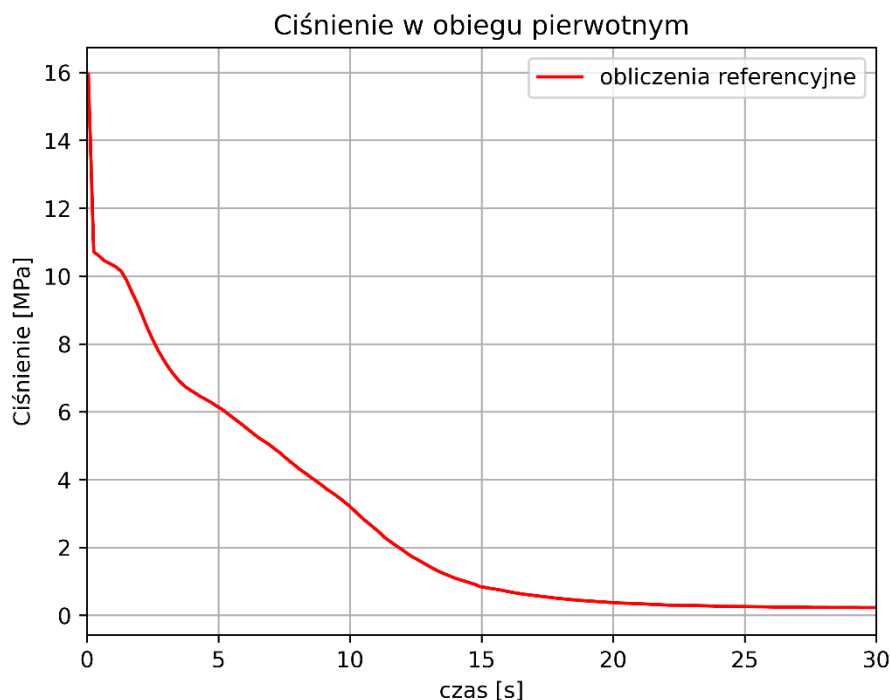
Rys. 6.9. Przebieg poziomu wody w rdzeniu dla obliczeń referencyjnych dla modelu elektrowni jądrowej



Rys. 6.10. Przebieg masowego natężenia przepływu przez rozerwanie dla obliczeń referencyjnych dla modelu elektrowni jądrowej

Na rysunku 6.10 przedstawiono przebieg masowego natężenia przepływu przez rozerwanie w ograniczeniu do trzydziestu pierwszych sekund, ponieważ później ustala się stała wartość wydatku przepływu. Przebieg ten przedstawia zsumowane wartości z obu zaworów symulujących obustronne rozerwanie. Maksymalne natężenie przepływu zaobserwowano po 1.31 sekundzie a jego wartość wynosiła 33 076 kg/s.

Na rysunku 6.11 przedstawiono przebieg ciśnienia w obiegu pierwotnym, w pętli bez rozerwania, przy ograniczeniu do trzydziestu pierwszych sekund, ponieważ później ustala się stała wartość ciśnienia, równa ciśnieniu panującemu w obudowie bezpieczeństwa. W pierwszej sekundzie ciśnienie gwałtownie spada do 10.5 MPa, następnie w ciągu 25 sekund spada do 0.25 MPa. Przebieg zmian parametru jest więc typowy dla awarii typu LBLOCA. Szybki spadek ciśnienia pozwala na wyzwolenie sygnałów uruchomienia wtrysku z systemów awaryjnego rdzenia w krótkim czasie po rozpoczęciu awarii. W pierwszej kolejności rozpoczyna się wtrysk ze zbiorników akumulatorów, których zawory zwrotne otwierają się, gdy ciśnienie spadnie poniżej 4.25 MPa. Następnie następuje wtrysk z HPIS oraz LPIS.



Rys. 6.11. Przebieg ciśnienia w obiegu pierwotnym (pętla bez rozerwania) dla obliczeń referencyjnych dla modelu elektrowni jądrowej

6.2.3. Obliczenia niepewności stosując uzyskane rozkłady gęstości prawdopodobieństwa parametrów (element szósty metodyki)

Ostatnim krokiem metodyki jest wykonanie obliczeń niepewności dla reaktora energetycznego w których zmieniane będą zarówno wewnętrzne parametry modelowe kodu oraz inne parametry określające między innymi warunki początkowe. Rozkłady gęstości prawdopodobieństwa dla parametrów modelowych określone zostały dla dwóch metod kwantyfikacji niepewności w [rozdziałach 5.1.4](#) oraz [5.2.4](#). W tabeli 6.7 przedstawiono natomiast parametry reprezentujące moc reaktora oraz warunki początkowe, wraz z ich zakresem zmienności oraz określeniem rozkładu. Dla wszystkich tych parametrów zastosowano rozkłady jednostajne, gdyż wszystkie wartości z zakresu zmienności są równie prawdopodobne. Trzy parametry reprezentują zmiany mocy reaktora, każdy z nich wprowadzany jest w formie mnożnika dla wartości referencyjnych. Trzy parametry reprezentujące warunki początkowe to poziom wody w stabilizatorze ciśnienia, prędkość obrotowa pompy cyrkulacyjnej obiegu pierwotnego oraz ciśnienie w akumulatorze.

Tabela 6.7. Zastosowane w analizie niepewności parametry reprezentujące moc reaktora oraz warunki początkowe wraz z ich zakresem zmian oraz rozkładem gęstości prawdopodobieństwa

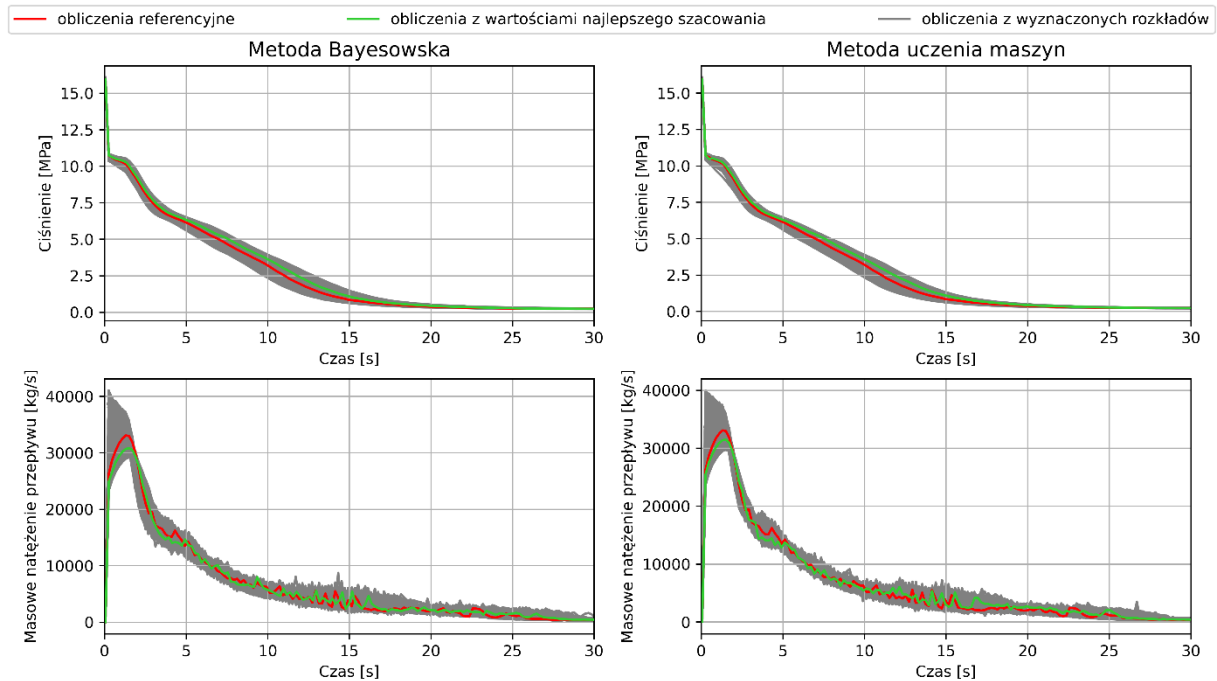
	Parametr	Zakres zmian	Rozkład
Parametry mocy reaktora	Moc reaktora	0.98 – 1.02	Jednostajny
	Mnożnik ciepła powyłączeniowego	0.98 – 1.02	Jednostajny
	Współczynnik nierównomierności rozkładu mocy w rdzeniu	0.98 – 1.02	Jednostajny
Warunki początkowe	Poziom wody w stabilizatorze ciśnienia	± 0.1 m	Jednostajny
	Prędkość obrotowa pompy cyrkulacyjnej	0.98 – 1.02	Jednostajny
	Ciśnienie w akumulatorze	± 0.1 MPa	Jednostajny

Obliczenia niepewności wykonano więc łącznie dla 15 parametrów stosując metodę propagacji niepewności wejściowych, opisaną w [rozdziale 2.1.1](#). Obliczenia niepewności wykonywano dla dwóch zestawów rozkładów dziewięciu parametrów modelowych kodu - uzyskanych stosując metodę Bayesowską oraz metodę uczenia maszyn. Z każdego zestawu tych parametrów losowano 500 wartości parametrów. Następnie losowano 500 wartości parametrów opisanych w tabeli 6.7. Wartości te wprowadzano do plików wsadowych kodu TRACE opisujących model reaktora typu PWR. Po przygotowaniu 500 plików wsadowych dla każdego z zestawu parametrów, wykonano łącznie 1000 obliczeń (po 500 dla każdej z metod na podstawie, której określono rozkłady parametrów). Obliczenia dla jednego przypadku trwały około 2 godzin, stąd przy wykorzystaniu sześciu rdzeni komputera całkowity czas obliczeń wyniósł dwa tygodnie. Dodatkowo dla każdej z metod wykonano obliczenia z wartościami najlepszego szacowania parametrów, określonymi w rozdziale piątym.

Zgodnie z formułą Wilksa, aby osiągnąć przedział tolerancji 95/95 dla 500 obliczeń można z kryterium akceptacji porównywać wartość siedemnastej największej wartości wielkości wyjściowej. Stosując statystykę porządkową [109] względem badanej wielkości wyjściowej (np. maksymalnej temperatury koszulki paliwowej) odrzuca się, więc 16 obliczeń dla których badana wielkość była największa, a siedemnastą porównuje z kryterium akceptacji.

Na rysunku 6.12 przedstawiono wyniki obliczeń walidacyjnych dla ciśnienia w obiegu pierwotnym oraz dla masowego natężenia przepływu. Podobnie jak na rysunkach 6.10 oraz 6.11 przebiegi przedstawiono do trzydziestej sekundy, ponieważ dalej ustalają się stałe wartości. Uzyskane z zastosowaniem obu metod zakresy niepewności są zbieżne ze sobą, dostrzegalne są jedynie minimalne różnice. Zakresy niepewności w przebiegu ciśnienia są wąskie, co wskazuje, że zastosowane rozkłady gęstości prawdopodobieństwa są poprawne.

Reaktor PWR obliczenia walidacyjne



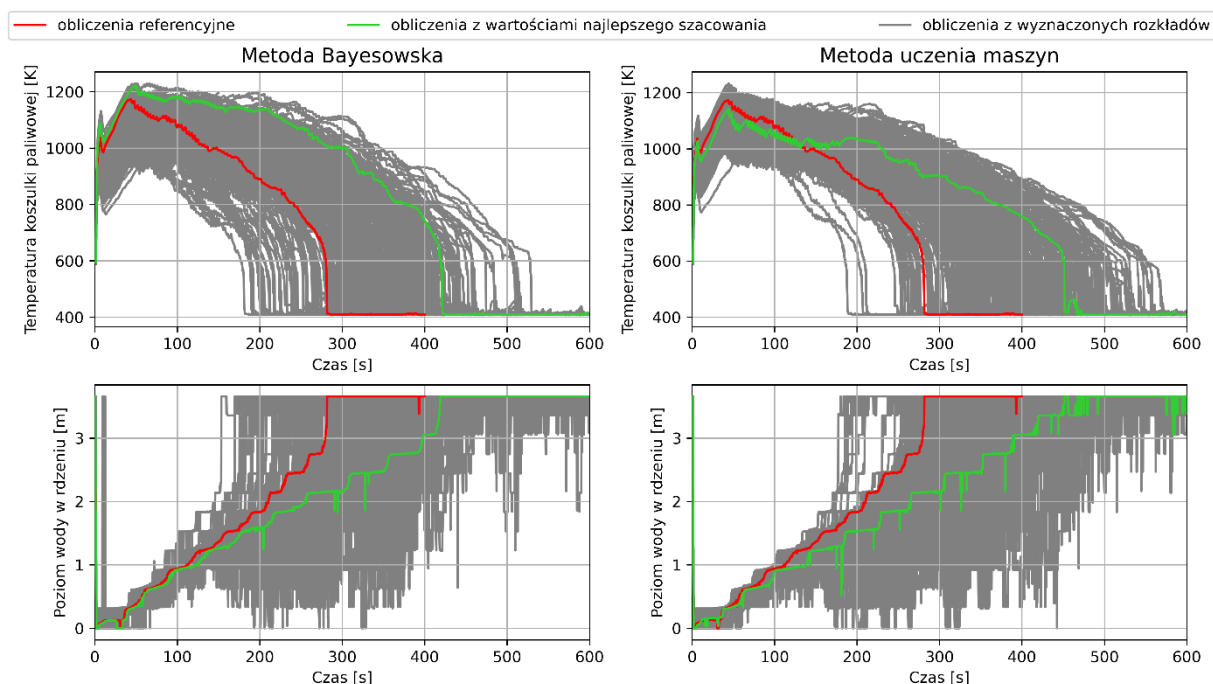
Rys. 6.12. Zakresy niepewności ciśnienia w obiegu oraz masowego natężenia przepływu w obliczeniach walidacyjnych dla reaktora energetycznego typu PWR stosując rozkłady parametrów wejściowych uzyskane z zastosowaniem dwóch opracowanych metod

W przypadku masowego natężenia przepływu obliczenia najlepszego szacowania są w obu przypadkach zbliżone do obliczeń referencyjnych. Zgodnie z wyznaczonymi wartościami współczynnika CHM12 mniejszymi od wartości domyślnej (równiej 1.0) maksymalna wartość wydatku przepływu podczas wypływu cieczy przechłodzonej jest mniejsza w obliczeniach najlepszego szacowania. Maksymalne wartości wydatku przepływu wyznaczone w obliczeniach walidacyjnych osiągają wartość 40 000 kg/s co stanowi wartość o 21% większą od maksymalnej wartości w obliczeniach referencyjnych. Dalej do 30 sekundy zakres niepewności wynosi $\pm 15\%$ wartości referencyjnej. Zakresy niepewności nie są więc bardzo szerokie, ale również nie są nierealistycznie wąskie.

Na rysunku 6.13 przedstawiono wyniki obliczeń walidacyjnych dla temperatury koszulki paliwowej oraz poziomu wody w rdzeniu. W przypadku temperatury koszulki paliwowej najistotniejsze jest, że górna wartość zakresu niepewności nie przekroczyła limitu równego 1473 Kelvinów. W przypadku obliczeń wykonanych dla rozkładów uzyskanych stosując metodę Bayesowską maksymalna temperatura koszulki paliwowej wynosiła 1205 K, natomiast w przypadku metody uczenia maszyn 1220 K. Dla metody uczenia maszyn uzyskano węższy zakres niepewności. Obliczenia te charakteryzują się, jednakże dłuższym czasem uzyskiwania

minimalnej wartości temperatury koszulki paliwowej. Jest to zgodne z przebiegiem obliczeń najlepszego szacowania dla metody uczenia maszyn. Zakresy niepewności dla poziomu wody w rdzeniu dla obu metod są szerokie. Jest to uzależnione od szerokiego zakresu zmian czasu, w którym temperatura koszulki paliwowej osiąga wartość minimalną.

Reaktor PWR obliczenia walidacyjne



Rys. 6.13. Zakresy niepewności temperatury koszulki paliwowej oraz poziomu wody w rdzeniu w obliczeniach walidacyjnych dla reaktora energetycznego typu PWR stosując rozkłady parametrów wejściowych uzyskane z zastosowaniem dwóch opracowanych metod

Podobnie jak dla obliczeń walidacyjnych wykonanych dla instalacji LOFT, tak również dla obliczeń niepewności dla reaktora energetycznego nieznacznie węższe zakresy niepewności badanych wielkości wyjściowych uzyskano dla parametrów określonych z zastosowaniem metody wstecznej kwantyfikacji niepewności opartej o uczenie maszyn. Uzyskane zakresy niepewności są szerokie, jednak pozwalają na pokrycie danych eksperymentalnych. Zgodnie ze wstępnymi wnioskami przedstawionymi w rozdziale 5.2.5 uzyskiwanie szerokich zakresów niepewności było spodziewane. Wynika to z metody wyznaczania rozkładów gęstości prawdopodobieństwa, których zakresy były szersze ze względu na złożenie rozkładów dla wielu testów o różnych warunkach początkowych i brzegowych w instalacjach eksperymentalnych. Zapewnia to większą uniwersalność uzyskanych rozkładów, czyli możliwość zastosowania ich przy zróżnicowanych warunkach początkowych w obiektach jądrowych.

7. Podsumowanie

W rozprawie zaprezentowano usprawnienie dotychczasowych metod wykonywania obliczeń najlepszego szacowania wraz z oceną niepewności poprzez zaproponowanie metodyki, która uwzględnia kwantyfikację niepewności parametrów modelowych kodów obliczeniowych stosowanych do wykonywania obliczeń cieplno-przepływowych w analizach deterministycznych obiektów jądrowych. Metodyka bazuje na połączeniu najlepszych praktyk wykonywania obliczeń niepewności przy wykorzystaniu dotychczasowo stosowanych metod, jednocześnie pozwalając na minimalizację niepewności wynikających z efektu użytkownika. Sześciokrokowa metodyka opisana została w rozdziale trzecim. Istotną częścią metodyki jest przeprowadzenie wstecznej kwantyfikacji niepewności parametrów modelowych. Parametry modelowe dostępne w kodach obliczeniowych najczęściej w formie mnożników, pozwalają na kalibrację wyników obliczeń. Pożądane jest, aby w systematyczny i oparty na danych eksperymentalnych sposób wyznaczyć dopuszczalne zakresy zmienności i rozkłady tych parametrów. W innym wypadku użytkownicy kodów mogą przyjmować dowolne wartości tych współczynników, co może prowadzić do wątpliwego jakościowo nadmiernego dopasowywania wyników do kryteriów akceptacji bądź danych pomiarowych. W celu określenia prawdziwych zakresów stosowności oraz rozkładów gęstości prawdopodobieństwa parametrów wykorzystuje się wsteczną kwantyfikację niepewności parametrów. W ramach rozprawy opracowano dwie metody obliczeniowe pozwalające na przeprowadzenie takiej kwantyfikacji. Pierwsza z metod bazuje na wnioskowaniu Bayesowskim wspieranym przez wykonywanie obliczeń Monte Carlo stosując łańcuch Markova, natomiast druga oparta jest o algorytm uczenia maszyn. Obie metody opisano w rozdziale czwartym. Implementację obu metod wykonano dla parametrów modelowych kodu obliczeniowego TRACE opisujących trzy wybrane zjawiska obserwowane podczas awarii typu LBLOCA: wpływ krytyczny, wymianę ciepła w rdzeniu oraz propagację frontu zalewania. Obliczenia wykonano dla dwóch instalacji eksperymentalnych, w których badano te zjawiska: Marviken oraz Flecht Seaset. Rezultatem wykonanej kwantyfikacji były rozkłady gęstości prawdopodobieństwa dziewięciu wybranych parametrów modelowych kodu TRACE. Rozkłady te określają możliwe wartości, jakie użytkownik może przyjąć dla badanych parametrów modelowych, aby osiągnąć wyniki zbieżne z danymi eksperymentalnymi przy zróżnicowanych warunkach początkowych oraz brzegowych. Rozdział piąty przedstawia w kompleksowy sposób implementację obu opracowanych metod obliczeniowych do problemu wstecznej kwantyfikacji niepewności dla instalacji eksperymentalnych Marviken i Flecht Seaset. Opracowana metodyka przewiduje również wykonanie walidacyjnych obliczeń niepewności z zastosowaniem uzyskanych

rozkładów gęstości prawdopodobieństwa parametrów modelowych, aby ocenić poprawność określenia tych rozkładów. Obliczenia walidacyjne wykonano dla instalacji eksperymentalnej LOFT, w której symulowano awarię typu LBLOCA oraz dla modelu reaktora energetycznego PWR. Wyniki obliczeń walidacyjnych przedstawiono w rozdziale szóstym.

W pracy zaproponowano nową metodykę oraz dwie metody obliczeniowe, z których jedna jest nowatorska i dotychczas zaprezentowana została w kontekście kwantyfikacji niepewności parametrów w analizach deterministycznych obiektów jądrowych jedynie przez autora rozprawy [101], [103]. Wdrożenie metodyki w stopniu zaprezentowanym w rozprawie również nie zostało dotychczas zaprezentowane w literaturze tematycznej.

7.1. Ocena metodyki oraz zaproponowanych metod

Motywacją rozprawy było usprawnienie obecnie stosowanej metody propagacji niepewności wejściowych, które pozwoli na ograniczenie wpływu użytkownika oraz zapewnienie większej uniwersalności oraz powtarzalności metody. Zaprezentowana metodyka oferuje systematyczną, uniwersalną oraz powtarzalną sposobność na określenie zakresów i rozkładów gęstości prawdopodobieństwa dla parametrów modelowych. Poprzez kwantyfikację niepewności tych parametrów z zastosowaniem danych eksperymentalnych redukowany jest wpływ użytkownika. Dodatkowo określenie zakresu stosowalności na podstawie danych eksperymentalnych pozwala na uzyskanie pewności, że nie będą stosowane wartości, które nie pozwalają na poprawne odwzorowanie warunków panujących w eksperymentach stosowanych przy wykonywaniu walidacji kodów obliczeniowych.

Opracowana metodyka odwzorowuje najlepsze współczesne praktyki wykonywania obliczeń niepewności zgodnie z rekomendacjami zawartymi w projekcie SAPIUM realizowanym pod egidą OECD. Pierwszy element metodyki, czyli specyfikacja problemu, jest realizowany każdorazowo w deterministycznych analizach bezpieczeństwa niezależnie od metody. Drugi element metodyki to przygotowanie odpowiedniej bazy eksperymentalnej adekwatnej do określonego problemu. Wymagane jest więc posiadanie dostępu do danych eksperymentalnych z różnych instalacji testowych, co można realizować poprzez dostęp do bazy danych OECD NEA bądź wykorzystując publicznie dostępne dane. W rozprawie zaprezentowano typowe podejście do określenia istotnych zjawisk zachodzących podczas awarii LBLOCA oraz określenia instalacji eksperymentalnych w których badano te zjawiska. Element trzeci metodyki to opracowanie modeli określonych w poprzednim kroku instalacji eksperymentalnych oraz poddanie ich procesowi kwalifikacji. W rozprawie zaprezentowano metody kwalifikacji modeli obiektów jądrowych, które mogą być uniwersalnie stosowane przy

wykonywaniu deterministycznych analiz bezpieczeństwa niezależnie od stosowanego podejścia (konserwatywnego bądź najlepszego szacowania). Na tym etapie określa się również, na podstawie globalnej analizy wrażliwości, najbardziej istotne parametry, które będą analizowane w kolejnym kroku. W rozprawie zaprezentowano metodę globalnej analizy wrażliwości opartą o analizę wariancji. Element czwarty metodyki to kwantyfikacja niepewności określonych w poprzednim kroku parametrów modelowych kodu obliczeniowego, poprzez zastosowanie opisanych metod wstecznej kwantyfikacji niepewności. Element czwarty jest więc elementem, w którym wykorzystano dwie opracowane w rozprawie metody obliczeniowe. Jest to obszar, nad którym prowadzone jest obecnie wiele badań oraz prób określenia najlepszych i najbardziej efektywnych metod. Elementy piąty i szósty to walidacja uzyskanych zakresów niepewności parametrów dla instalacji eksperymentalnej oraz końcowe obliczenia niepewności dla reaktora energetycznego. Metodyka stanowi więc kompleksowe podejście do tematyki wykonywania obliczeń najlepszego szacowania wraz z oceną niepewności dla obiektów jądrowych. Metodyka pozwala na uwzględnienie wielu źródeł niepewności, w tym określenie niepewności parametrów modelowych.

Znacząca część rozprawy poświęcona była prezentacji opracowanych metod wstecznej kwantyfikacji niepewności parametrów. Pierwsza z metod, oparta o wnioskowanie Bayesowskie oraz wykonywanie obliczeń Monte Carlo stosując łańcuch Markova, posiada jasno określone podstawy matematyczne. O ile samo twierdzenie Bayesa jest intuicyjne, to jego implementacja do praktycznych problemów określenia rozkładów a’posteriori jest już problemem nietrywialnym, stąd najczęściej wykorzystuje się metody Monte Carlo do losowania wartości z przestrzeni próbkowania. W rozprawie wykorzystano algorytm próbkowania Metropolis-Hastingsa, który pozwala na uproszczenie procesu porównywania rozwiązań funkcji wiarygodności uzyskanych po wykonaniu obliczeń z zaproponowanymi (poprzez próbkowanie stosując łańcuch Markova) wartościami parametrów. Algorytm wymaga określenia przez użytkownika pierwszego oszacowania wartości początkowych parametrów, rozkładu a’priori parametrów oraz rozkładu propozycji, czyli rozkładu, z którego będą losowane propozycje wartości parametrów generowane przez łańcuch Markova. O ile jako wartości pierwszego oszacowania można przyjmować wartości domyślne parametrów to określenie rozkładu propozycji wraz z jego parametryzacją statystyczną (np. wartość średnia i odchylenie standardowe w przypadku rozkładu normalnego), będzie miało znaczący wpływ na uzyskiwane wyniki kwantyfikacji niepewności, stąd jest to element wciąż podatny na wpływ użytkownika. W rozprawie wykorzystywano rozkład normalny o stałym odchyleniu

standardowym oraz wartości średniej przyjmującej wartość parametru w poprzednim kroku, co pozwoliło na kolejne uproszczenie procesu porównywania zaproponowanych w łańcuchu Markova wartości parametrów. Zaletą stosowania tej metody jest możliwość uwzględnienia w obliczeniach rozbieżności modelowej. Dodatkowo poszukiwane rozkłady a’posteriori są od razu uzyskiwane w wyniku zastosowania metody bez konieczności dodatkowej obróbki danych.

Metoda oparta o algorytm uczenia maszyn jest oryginalnym wkładem w metody kwantyfikacji niepewności parametrów w deterministycznych analizach obiektów jądrowych. Algorytm nadzorowanego uczenia maszyn zaimplementowano do oceny miary rozbieżności między danymi eksperymentalnymi a obliczeniami wykonywanymi z zastosowaniem kodu obliczeniowego. Metoda ta posiada wciąż pewną wrażliwość na efekt użytkownika, jednakże opisane w rozprawie metody tworzenia klas, zbioru uczącego oraz określenia cech mogą stanowić wytyczne pozwalające na minimalizację wpływu użytkownika. Metoda wymaga większej liczby kroków pośrednich niż metoda Bayesowska. Konieczne jest przygotowanie zbioru uczącego, określenie klas opisujących poziom zbieżności między danymi eksperymentalnymi a wynikami obliczeń oraz przypisanie ręczne klas do obliczeń ze zbioru uczącego. Kolejne kroki są już bardziej zautomatyzowane i polegają na opisanu obliczeń zbioru uczącego wartościami określonych cech i wyuczeniu się przez algorytm relacji między cechami a klasami dla zbioru uczącego. Następnie wykonywane są obliczenia dla zbioru testowego, a algorytm przypisuje automatycznie klasy do obliczeń na podstawie wyznaczenia ich cech oraz znajomości relacji między cechami a klasami. Wartości parametrów, które pozwoliły na uzyskanie wyników przypisanych do klasy A (dobra zgodność między danymi eksperymentalnymi a obliczeniami) stanowią poszukiwane rozkłady a’posteriori. Mimo wielu koniecznych do implementacji kroków metoda jest intuicyjna. Główną różnicą w stosunku do metody Bayesowskiej jest przestrzeń próbkowania spośród której losowane są parametry. W metodzie Bayesowskiej poprzez zastosowanie łańcucha Markova próbuje się z przestrzeni rozkładu a’posteriori, a w przypadku metody opartej o uczenie maszyn z przestrzeni rozkładu a’priori. Powoduje to konieczność wykonywania większej liczby obliczeń w metodzie opartej o uczenie maszyn. W przypadku obliczeń dla instalacji Marviken stosując metodę Bayesowską zaakceptowano 33% propozycji, natomiast w przypadku metody opartej o algorytm uczenia maszyn było to 7% obliczeń zaklasyfikowanych do klasy A. W przypadku obliczeń dla instalacji Flecht Seaset było to odpowiednio 56% dla temperatury koszulki paliwowej i 70% dla propagacji frontu zalewania dla metody Bayesowskiej i 14% oraz 19% dla metody opartej

o algorytm uczenia maszyn. O ile liczba obliczeń jest zdecydowanie większa dla metody opartej o uczenie maszyn to sam czas wykonywania obliczeń może być krótszy. W przypadku metody Bayesowskiej wartości parametrów losowane są każdorazowo z rozkładu propozycji zależnego od poprzedniego kroku, stąd plik wsadowy do obliczeń przygotowywany jest bezpośrednio po ukończeniu poprzednich obliczeń. Nie ma więc możliwości wykonywania wielu obliczeń równolegle. W przypadku metody opartej o uczenie maszyn wszystkie pliki wsadowe tworzone są w jednym momencie poprzez proste losowe próbkowanie z rozkładów a’piori parametrów. Pozwala to na wykonywanie wielu obliczeń jednocześnie, a ich liczba jest jedynie ograniczona liczbą rdzeni i wątków maszyny roboczej.

Podczas wykonywania kwantyfikacji niepewności parametrów modelowych wyniki obliczeń porównywane były z danymi eksperymentalnymi. W ramach porównania poprawności uzyskiwanych wyników dokonywano zarówno analizy jakościowej, jak i ilościowej. W przypadku obliczeń wykonywanych dla zjawiska przepływu krytycznego lepsze dopasowanie do danych eksperymentalnych uzyskiwano stosując metodę opartą o algorytm uczenia maszyn. Podobne obserwacje poczyniono dla obliczeń wykonywanych dla temperatury koszulki paliwowej oraz propagacji frontu zalewania. Uzyskane obiema metodami rozkłady gęstości prawdopodobieństwa parametrów zweryfikowano poprzez wykonanie obliczeń sprawdzających dla instalacji Marviken i Flecht Seaset oraz instalacji LOFT. W większości przypadków wyniki obliczeń sprawdzających pozwoliły na pokrycie zakresu błędu pomiarowego danych eksperymentalnych. Węższe zakresy uzyskano stosując metodę opartą o uczenie maszyn, jednak rozbieżności między wynikami uzyskiwanymi stosując obie metody nie były znaczące.

Uzyskane zakresy niepewności były szerokie, co wynikało z faktu, że proponowane rozkłady gęstości prawdopodobieństwa były kompromisem pomiędzy rozkładami uzyskiwanymi dla obliczeń dla testów o różnych warunkach początkowych. Jest to trudny do jednoznacznego rozwiązania problem poszukiwania rozkładu uniwersalnego, jednak mniej dokładnego lub dającego dokładne wyniki, które mogą być jednak poprawne jedynie w wąskim zakresie stosowalności. W rozprawie zdecydowano się poszukiwać rozkładów uniwersalnych. W związku z tym zakresy niepewności wielkości wyjściowych były szersze. Dodatkową konsekwencją jest założenie, że wartości parametrów wejściowych poza wyznaczonymi rozkładami gęstości prawdopodobieństwa nie powinny być stosowane.

Kolejnym problemem był fakt, że o ile konkretne zestawy parametrów z uzyskanych w wyniku wstecznej kwantyfikacji zakresów niepewności pozwalały na uzyskanie dobrego dopasowania

do danych eksperymentalnych, co skutkowało wąskimi przedziałami niepewności obserwowanych wielkości wyjściowych, to już próbkowanie losowe z tych zakresów nie daje aż tak wąskich przedziałów niepewności. W obrębie wyznaczonych równocześnie dla wielu parametrów rozkładów, istnieją więc kombinacje dające dobre dopasowanie do danych eksperymentalnych jak i kombinacje, dla których wyniki są bardziej rozbieżne.

7.2. Wnioski i możliwości kontynuacji badań

Spśród dwóch zaproponowanych metod wstecznej kwantyfikacji niepewności węższe zakresy niepewności wielkości wyjściowych opisujących badane zjawiska uzyskiwano dla metody opartej o uczenie maszyn. Metoda ta pozwala również na wyszczególnienie najbardziej istotnych cech do oceny ilościowej wyników, co jest istotne między innymi przy badaniu parametrów takich jak temperatura koszulki paliwowej dla których określone są liczbowe kryteria akceptacji. Proponowane jest więc stosowanie rozkładów dziewięciu parametrów modelowych uzyskanych metodą uczenia maszyn i zaprezentowanych w rozdziałach [5.1.4.2](#) oraz [5.2.4.2](#).

Przyszłe badania powinny dotyczyć rozwoju obu metod wstecznej kwantyfikacji niepewności. Dla metody opartej o wnioskowanie Bayesowskie w pierwszej kolejności należałoby zbadać możliwość zastosowania metamodeli do zastąpienia wykonywania obliczeń z kodem cieplno-przepływowym. Należałoby więc opracować taki metamodel i porównać jego dokładność w odniesieniu do obliczeń wykonywanych przy zastosowaniu kodu TRACE. Gdyby metamodel gwarantował dobrą dokładność w odwzorowaniu relacji między zmianami parametrów wejściowych a zmianami badanych wielkości wyjściowych, wówczas znacząco skróciłoby to czas obliczeń.

Dla metody opartej o uczenie maszyn przyszłe badania powinny dotyczyć przede wszystkim ustalenia klarownych reguł wyboru testów do zbioru uczącego oraz zbioru testowego. Dodatkowo zaobserwowano, że określanie jednego zbioru uczącego dla całego zbioru testowego może powodować przypisywanie niskiej liczby obliczeń do klasy A. W związku z tym przyszłe badania powinny skupić się na możliwości podziału bazy testów wykonanych w jednej instalacji, tak aby dla testów o podobnych warunkach początkowych i brzegowych utworzyć podgrupy w których jeden test znalazłby się w zbiorze uczącym, a drugi w zbiorze testowym. Dla każdej podgrupy algorytm trenowany byłby oddzielnie. W drugiej kolejności można usprawnić proces przypisywania klas do obliczeń ze zbioru uczącego poprzez bardziej restrykcyjne reguły przypisywania oraz eliminację podobnych rezultatów ze zbioru.

Dodatkowo niezależnie od wyboru metody wstecznej kwantyfikacji niepewności parametrów powinno wykonywać się dokładane lokalne badania wrażliwości dla wszystkich wybranych parametrów, aby lepiej określać rozkłady oraz zakresy zmienności a’priori. Pozwoli to ograniczyć liczbę wykonywanych obliczeń, które będą znacząco rozbieżne od danych eksperymentalnych.

Wykonanie badań and rozwojem metod pozwoli również na rozszerzenie bazy eksperymentalnej dla awarii LBLOCA analizowanej w ramach zaproponowanej metodyki i wykonanie obliczeń dla kolejnych zjawisk. Można będzie wówczas powtórzyć obliczenia dla instalacji LOFT oraz modelu elektrowni jądrowej z większą liczbą parametrów, bądź zaktualizowanymi rozkładami. Opracowana metodyka, wraz z wykonanymi obliczeniami oraz modelami może, więc stanowić podstawę do dalszego rozwoju bazy eksperymentalnej, parametrów oraz wykonywania kolejnych obliczeń dla innych modeli reaktorów energetycznych. Ostatecznym celem będzie więc wdrożenie tej metodyki dla obliczeń najlepszego szacowania wraz z oceną niepewności wykonywanych dla reaktora AP1000 wybranego przez Polski Rząd jako technologia, w której budowana będzie pierwsza elektrownia jądrowa w Polsce. Metodyka pozwoli na przeprowadzenie przez Państwową Agencję Atomistyki sprawdzających obliczeń BEPU dla wybranych scenariuszy awaryjnych.

8. Referencje

- [1] IAEA, “GSR 4: Safety assessment for facilities and activities,” 2016.
- [2] IAEA, “Deterministic Safety Analysis for Nuclear Power Plants - IAEA SSG-2 Rev. 1,” 2019.
- [3] G. Petrangeli, *Nuclear Safety*. Elsevier, 2006.
- [4] *Ustawa z dnia 29 listopada 2000 r. - Prawo atomowe*. 2019, p. Dz.U.2021.0.1941.
- [5] *Rozporządzenie Rady Ministrów z dnia 31 sierpnia 2012 r. w sprawie wymagań bezpieczeństwa jądrowego i ochrony radiologicznej, jakie ma uwzględniać projekt obiektu jądrowego*. p. Dziennik ustaw z 2012 r. pozycja 1048.
- [6] *Rozporządzenie Rady Ministrów z dnia 31 sierpnia 2012 r. w sprawie zakresu i sposobu przeprowadzania analiz bezpieczeństwa przeprowadzanych przed wystąpieniem z wnioskiem o wydanie zezwolenia na budowę obiektu jądrowego, oraz zakresu wstępnego raportu bez.* p. Dziennik ustaw z 2012 r. pozycja 1043.
- [7] F. D’Auria, A. Bousbia-salah, A. Petruzzi, and a. D. Nevo, “State of the Art in Using Best Estimate Calculation Tools in Nuclear Technology,” *Nucl. Eng. Technol.*, vol. 38, no. 1, pp. 11–32, 2006.
- [8] F. D’Auria, *Thermal-Hydraulics of Water Cooled Nuclear Reactors*. Elsevier, 2017.
- [9] IAEA, “Derivation of the source term and analysis of the radiological consequences for the design basis accidents at research reactor,” *Saf. Reports Ser. No. 53*, p. 178, 2008.
- [10] N. W. Porter and V. A. Mousseau, “Understanding aleatory and epistemic parameter uncertainty in statistical models,” in *Best Estimate Plus Uncertainty International Conference, Italy*, 2020.
- [11] S. Lutsanych, “Improving Best Estimate approaches with uncertainty quantification in nuclear reactor thermal-hydraulics,” UNIVERSITÀ DI PISA, 2016.
- [12] IAEA, “Best Estimate Safety Analysis for Nuclear Power Plants: Uncertainty Evaluation; Safety Reports Series 52,” 2008.
- [13] A. Petruzzi and F. D’Auria, “Approaches, relevant topics, and internal method for uncertainty evaluation in predictions of thermal-hydraulic system codes,” *Sci. Technol. Nucl. Install.*, vol. 2008, 2008.
- [14] OECD-NEA, “CSNI Integral Test Facility Validation Matrix For the assessment of Thermal-Hydraulic codes for LWR LOCA and Transients.”
- [15] OECD-NEA, “Separate effects test matrix for thermal-hydraulic code validation.pdf.” 1993.
- [16] J. Baccou *et al.*, “Development of good practice guidance for quantification of thermal-hydraulic code model input uncertainty,” *Nucl. Eng. Des.*, vol. 354, no. June, p. 110173, 2019.
- [17] B. Boyack *et al.*, “Quantifying Reactor Safety Margins: Application of Code Scaling, Applicability, and Uncertainty Evaluation Methodology to a Large-Break, Loss-of-Coolant Accident,” *NUREG/CR-5249 EGG-2552 ADAMS Access. Number ML030380503*, 1989.
- [18] E. Hofer, *Probabilistische Unsicherheitsanalyse von Ergebnissen umfangreicher Rechenmodelle: Abschlußbericht*. Ges. für Anlagen- und Reaktorsicherheit (GRS) mbH, 1993.
- [19] V. Forge, A., Pochard, R., Gleaser H. Hobbhahn, W. Hofer, E., Technedorff, “Review study on uncertainty methods for thermal-hydraulic computer codes,” 1994.
- [20] R. P. Martin and A. Petruzzi, “Progress in international best estimate plus uncertainty analysis methodologies,” *Nucl. Eng. Des.*, vol. 374, no. January, p. 111033, 2021.

- [21] F. D'Auria, "Best Estimate Plus Uncertainty (BEPU): Status and perspectives," *Nucl. Eng. Des.*, vol. 352, no. July, p. 110190, 2019.
- [22] F. D'Auria, C. Camargo, and O. Mazzantini, "The Best Estimate Plus Uncertainty (BEPU) approach in licensing of current nuclear reactors," *Nucl. Eng. Des.*, vol. 248, pp. 317–328, 2012.
- [23] N. Aksan, "An overview on thermal-hydraulic phenomena for water cooled nuclear reactors; part I: SETs, and ITFs of PWRs, BWRs, VVERs," *Nucl. Eng. Des.*, vol. 354, no. August, p. 110212, 2019.
- [24] S. S. Wilks, "Determination of Sample Sizes for Setting Tolerance Limits," *Ann. Math. Stat.*, vol. 12, no. 1, pp. 91–96, 1941.
- [25] F. D'auria, N. Debrecin, and G. M. Galassi, "Outline of the Uncertainty Methodology Based on Accuracy Extrapolation," *Nucl. Technol.*, vol. 109, no. 1, pp. 21–38, 1995.
- [26] R. Mendizábal, "Bayesian perspective in BEPU licensing analysis," *Nucl. Eng. Des.*, vol. 355, no. September, p. 110310, 2019.
- [27] X. Wu, T. Kozłowski, H. Meidani, and K. Shirvan, "Inverse uncertainty quantification using the modular Bayesian approach based on Gaussian process, Part 1: Theory," *Nucl. Eng. Des.*, vol. 335, no. June, pp. 339–355, 2018.
- [28] X. Wu, T. Kozłowski, H. Meidani, and K. Shirvan, "Inverse uncertainty quantification using the modular Bayesian approach based on Gaussian process, Part 2: Application to Trace," *Nucl. Eng. Des.*, vol. 335, no. June, pp. 417–431, 2018.
- [29] J. P. Yurko, "Uncertainty Quantification in Safety Codes Using a Bayesian Approach with Data from Separate and Integral Effect Tests," p. 366, 2014.
- [30] G. Roma *et al.*, "A Bayesian framework of inverse uncertainty quantification with principal component analysis and Kriging for the reliability analysis of passive safety systems," *Nucl. Eng. Des.*, vol. 379, no. December 2020, p. 111230, 2021.
- [31] Y. Liu, N. T. Dinh, R. C. Smith, and X. Sun, "Uncertainty quantification of two-phase flow and boiling heat transfer simulations through a data-driven modular Bayesian approach," *Int. J. Heat Mass Transf.*, vol. 138, pp. 1096–1116, 2019.
- [32] J. P. C. Kleijnen, "Kriging metamodeling in simulation: A review," *Eur. J. Oper. Res.*, vol. 192, no. 3, pp. 707–716, 2009.
- [33] A. Crécy, "Determination of the uncertainties of the constitutive relationships of the CATHARE 2 code," *M&C*, vol. 3, 2001.
- [34] I. Lee, D. Oh, Y. Bang, and Y. Kim, "Effect of critical flow model in MARS-KS code on uncertainty quantification of large break Loss of Coolant Accident (LBLOCA)," *Nucl. Eng. Technol.*, vol. 52, no. 4, pp. 755–763, 2019.
- [35] A. Kovtonyuk, S. Lutsanych, F. Moretti, and F. D'Auria, "Development and assessment of a method for evaluating uncertainty of input parameters," *Nucl. Eng. Des.*, vol. 321, pp. 219–229, 2017.
- [36] R. Mendizabál Sanz, E. de Alfonso, J. Freixa, and F. Reventós, "Post-BEMUSE Reflood Model Input Uncertainty Methods (PREMIUM) Benchmark: Final Report," no. August, 2017.
- [37] T. Skorek *et al.*, "Quantification of the uncertainty of the physical models in the system thermal-hydraulic codes – PREMIUM benchmark," *Nucl. Eng. Des.*, vol. 354, no. May, p. 110199, 2019.
- [38] F. D'Auria and W. Giannotti, "Development of a Code with the Capability of Internal

- Assessment of Uncertainty,” *Nucl. Technol.*, vol. 131, no. 2, pp. 159–196, 2000.
- [39] A. Petruzzi, F. D’Auria, W. Giannotti, and K. Ivanov, “Methodology of Internal Assessment of Uncertainty and Extension to Neutron Kinetics/Thermal-Hydraulics Coupled Codes,” *Nucl. Sci. Eng.*, vol. 149, no. 2, pp. 211–236, 2005.
 - [40] H. (GRS) Glaeser, P. (CEA) Bazin, J. (IRSN) Baccou, E. (IRSN) Chojnacki, and S. (IRSN) Destercke, “BEMUSE Phase VI Report Status report on the area, classification of the methods, conclusions and recommendations,” *Csni*, no. March, 2011.
 - [41] D. J. Diamond, “Experience using Phenomena Identification and Ranking Technique (PIRT) for nuclear analysis,” *PHYSOR-2006 - Am. Nucl. Soc. Top. Meet. React. Phys.*, vol. 2006, 2006.
 - [42] A. (CEA) de Crécy and P. (CEA) Bazin, “BEMUSE Phase III Report: Uncertainty and Sensitivity Analysis of the LOFT L2-5 Test,” *Bemuse*, no. May, 2007.
 - [43] T. Hollands, S. Buchholz, and A. Wielenberg, “Validation of the AC2 codes ATHLET and ATHLET-CD,” *Kerntechnik*, vol. 84, no. 5, pp. 397–405, 2019.
 - [44] F. (Upc) Reventós, M. (Upc) Pérez, and L. (Upc) Batet, “BEMUSE Phase V Report: Uncertainty and Sensitivity Analysis of a LB-LOCA in ZION Nuclear Power Plant,” *Bemuse*, 2009.
 - [45] F. Reventós, E. De Alfonso, and R. Mendizábal, “PREMIUM, a benchmark on the quantification of the uncertainty of the physical models in the system thermal-hydraulic codes: methodologies and data review - csni-r2016-9.pdf,” *CSNI Rep.*, no. April, 2016.
 - [46] S. T. Andrea Saltelli, Marco Ratto, Terry Andres, Francesca Campolongo, Jessica Cariboni, Debora Gatelli, Michaela Saisana, *Global Sensitivity Analysis: The Primer*. John Wiley & Sons, Incorporated, 2008.
 - [47] J. Zhang, A. Kovtonyuk, and C. Schneidesch, “Towards a graded application of best estimate plus uncertainty methodology for non-LOCA transient analysis,” *Nucl. Eng. Des.*, vol. 354, no. July, p. 110189, 2019.
 - [48] N. Aksan, F. D’Auria, and H. Glaeser, “Thermal-hydraulic phenomena for water cooled nuclear reactors,” *Nucl. Eng. Des.*, vol. 330, no. January, pp. 166–186, 2018.
 - [49] N. Systems, A. Operations, I. S. Laboratories, I. Falls, and U. S. N. R. Commission, “Relap5/Mod3.3 Code Manual Volume Viii: Programmers Manual,” *Control*, vol. VIII, no. October, 2010.
 - [50] U.S.NRC, “TRACE V5.0 Patch 5 USER’S MANUAL Volume 1: Input Specification,” vol. 1, p. 636, 2017.
 - [51] U.S.NRC, “TRACE V5.0 Patch 5 Theory Manual -- Field Equations, Solution Methods, and Physical Models,” 2017.
 - [52] R. Préa, P. Fillion, L. Matteo, G. Mauger, and A. Mekkas, “CATHARE-3 V2.1: the new industrial version of the CATHARE code,” 2020.
 - [53] D. Bestion, F. D’Auria, P. Lien, and H. Nakamura, “A state-of-the-art report mal- hydraulics applications to nuclear reactor safety and design,” 2017.
 - [54] C. Frepoli, “An overview of westinghouse realistic large break Loca evaluation model,” *Sci. Technol. Nucl. Install.*, vol. 2008, 2008.
 - [55] R. P. Martin and L. D. O’Dell, “AREVA’s realistic large break LOCA analysis methodology,” *Nucl. Eng. Des.*, vol. 235, no. 16, pp. 1713–1725, 2005.

- [56] S. Suehiro, J. Sugimoto, A. Hidaka, H. Okada, S. Mizokami, and K. Okamoto, "Development of the source term PIRT based on findings during Fukushima Daiichi NPPs accident," *Nucl. Eng. Des.*, vol. 286, pp. 163–174, 2015.
- [57] S. Gouénard, J.-P. Hélot, M. Biçakli, and E. Royer, "PIRT and V&V associated to an uncertainty quantification methodology for fluid thermal mixing," *Nucl. Eng. Des.*, vol. 366, p. 110736, 2020.
- [58] R. P. Martin, "Quantifying Phenomenological Importance in Best-Estimate Plus Uncertainty Analyses," *Nucl. Technol.*, vol. 175, no. 3, pp. 652–662, 2011.
- [59] J. Yurko and J. Buongiorno, "Quantitative Phenomena Identification and Ranking Table (QPIRT) for Reactor Safety Analysis," *Trans. Am. Nucl. Soc.*, vol. 104, 2014.
- [60] N. Aksan, "An overview on thermal-hydraulic phenomena for water cooled nuclear reactors; part II: ALWRs and SCWRs," *Nucl. Eng. Des.*, vol. 354, no. August, p. 110214, 2019.
- [61] Marviken Project Team, "The Marviken Full Scale Critical Flow Tests Report, 'Description of the Test Facility. MXC-101,'" 1979.
- [62] Y. Mu, X. Liu, and L. Wang, "A Pearson's correlation coefficient based decision tree and its parallel implementation," *Inf. Sci. (Ny)*, vol. 435, pp. 40–58, 2018.
- [63] J. van Doorn, A. Ly, M. Marsman, and E.-J. Wagenmakers, "Bayesian rank-based hypothesis testing for the rank sum test, the signed rank test, and Spearman's ρ ," *J. Appl. Stat.*, vol. 47, no. 16, pp. 2984–3006, 2020.
- [64] P. M. Stano and M. Spirzewski, "Modular global uncertainty analysis of event-driven indicators of system's availability," *Saf. Reliab. - Safe Soc. a Chang. World - Proc. 28th Int. Eur. Saf. Reliab. Conf. ESREL 2018*, no. June, pp. 2635–2644, 2018.
- [65] P. M. Stano and M. Spirzewski, "Bayesian Integrated Global Uncertainty & Sensitivity Analysis of System's Availability," 2019.
- [66] M. Spirzewski, H. Anglart, and P. M. Stano, "Uncertainty and sensitivity analysis of a phenomenological dryout model implemented in DARIA system code," *Nucl. Eng. Des.*, vol. 355, no. August, p. 110281, 2019.
- [67] Z. E. Marcinkowska, K. M. Pytel, and A. Frydrysiak, "Reactivity costs in MARIA reactor," *Ann. Nucl. Energy*, vol. 99, pp. 366–371, 2017.
- [68] M. Spirzewski, P. Domitr, and P. Darnowski, "Global uncertainty and sensitivity analysis of MELCOR and TRACE critical flow models against MARVIKEN tests," *Nucl. Eng. Des.*, vol. 378, no. 2021, p. 111150, 2021.
- [69] SNL, "MELCOR 2.1 Computer Code Manual - Code Assessments," 2016.
- [70] I. M. Sobol, "On quasi-Monte Carlo integrations," *Math. Comput. Simul.*, vol. 47, no. 2, pp. 103–112, 1998.
- [71] A. Saltelli, P. Annoni, I. Azzini, F. Campolongo, M. Ratto, and S. Tarantola, "Variance based sensitivity analysis of model output. Design and estimator for the total sensitivity index," *Comput. Phys. Commun.*, vol. 181, no. 2, pp. 259–270, 2010.
- [72] A. Saltelli, "Making best use of model evaluations to compute sensitivity indices," *Comput. Phys. Commun.*, vol. 145, no. 2, pp. 280–297, 2002.
- [73] W. Usher *et al.*, "SALib/SALib: v1.3.13." Zenodo, 2021.
- [74] "Anaconda Software Distribution," *Anaconda Documentation*. Anaconda Inc., 2020.
- [75] P. Gregory, *Bayesian Logical Data Analysis for the Physical Sciences: A Comparative*

Approach with Mathematica® Support. Cambridge University Press, 2005.

- [76] K. Campbell, “Statistical calibration of computer simulations,” *Reliab. Eng. Syst. Saf.*, vol. 91, no. 10, pp. 1358–1363, 2006.
- [77] J. Brynjarsdóttir and A. O’hagan, “Learning about physical parameters: The importance of model discrepancy,” *Inverse Probl.*, vol. 30, no. 11, 2014.
- [78] X. Wu, K. Shirvan, and T. Kozlowski, “Demonstration of the relationship between sensitivity and identifiability for inverse uncertainty quantification,” *J. Comput. Phys.*, vol. 396, pp. 12–30, 2019.
- [79] M. C. Kennedy and A. O’Hagan, “Bayesian calibration of computer models,” *J. R. Stat. Soc. Ser. B Stat. Methodol.*, vol. 63, no. 3, pp. 425–464, 2001.
- [80] J. S. Speagle, “A Conceptual Introduction to Markov Chain Monte Carlo Methods,” *arXiv*, 2019.
- [81] N. Metropolis, A. W. Rosenbluth, M. N. Rosenbluth, A. H. Teller, and E. Teller, “Equation of State Calculations by Fast Computing Machines,” *J. Chem. Phys.*, vol. 21, no. 6, pp. 1087–1092, 1953.
- [82] W. K. Hastings, “Monte Carlo sampling methods using Markov chains and their applications,” *Biometrika*, vol. 57, pp. 97–109, 1970.
- [83] P. Grechanuk, M. E. Rising, and T. S. Palmer, “Using Machine Learning Methods to Predict Bias in Nuclear Criticality Safety,” *J. Comput. Theor. Transp.*, vol. 47, no. 4–6, pp. 552–565, 2018.
- [84] M. I. Radaideh, D. Price, and T. Kozlowski, “Modeling Nuclear Data Uncertainties Using Deep Neural Networks,” *EPJ Web Conf.*, vol. 247, p. 15016, 2021.
- [85] G. M. Hobold and A. K. da Silva, “Automatic detection of the onset of film boiling using convolutional neural networks and Bayesian statistics,” *Int. J. Heat Mass Transf.*, vol. 134, pp. 262–270, 2019.
- [86] Y. Liu, N. Dinh, Y. Sato, and B. Niceno, “Data-driven modeling for boiling heat transfer: Using deep neural networks and high-fidelity simulation results,” *Appl. Therm. Eng.*, vol. 144, pp. 305–320, 2018.
- [87] K. Y. Chung, “A machine learning strategy with restricted sliding windows for real-time assessment of accident conditions in nuclear power plants,” *Nucl. Eng. Des.*, vol. 378, p. 111140, 2021.
- [88] W. Homenda and W. Pedrycz, *Pattern Recognition : A Quality of Data Perspective*,. John Wiley & Sons, Incorporated, 2018.
- [89] J. Freixa, E. de Alfonso, and F. Reventós, “Testing methodologies for quantifying physical models uncertainties. A comparative exercise using CIRCE and IPREM (FFTBM),” *Nucl. Eng. Des.*, vol. 305, pp. 653–665, 2016.
- [90] L. Breiman, “Random Forests,” *Mach. Learn.*, vol. 45, no. 1, pp. 5–32, 2001.
- [91] F. Pedregosa *et al.*, “Scikit-learn: Machine Learning in {P}ython,” *J. Mach. Learn. Res.*, vol. 12, pp. 2825–2830, 2011.
- [92] W. J. Tong L.S., *Thermal Analysis of Pressurized Water Reactors*, vol. 56, no. 2. American Nuclear Society, 1982.
- [93] N. E. Todreas and M. S. Kazimi, *Nuclear Systems I, Thermal Hydrualic Fundamentals*, vol. 20, no. 6. 1992.

- [94] M. Lanfredini *et al.*, “Critical flow prediction by system codes – Recent analyses made within the FONESYS network,” *Nucl. Eng. Des.*, vol. 366, no. February, p. 110731, 2020.
- [95] R. E. Henry and H. K. Fauske, “TWO-PHASE CRITICAL FLOW OF ONE-COMPONENT MIXTURES IN NOZZLES, ORIFICES, AND SHORT TUBES,” *J. Heat Transf.* 93 No. 2, 179-87(May 1971)., 1971.
- [96] J. A. Trapp and V. H. Ransom, “A choked-flow calculation criterion for nonhomogeneous, nonequilibrium, two-phase flows,” *Int. J. Multiph. Flow*, vol. 8, no. 6, pp. 669–681, 1982.
- [97] F. J. Moody, “Maximum Flow Rate of a Single Component, Two-Phase Mixture,” *J. Heat Transfer*, vol. 87, no. 1, pp. 134–141, 1965.
- [98] M. Alamgir and J. H. Lienhard, “Correlation of Pressure Undershoot During Hot-Water Depressurization,” *J. Heat Transfer*, vol. 103, no. 1, pp. 52–55, 1981.
- [99] U. NRC, “TRACE Appendix B : Separate Effects Tests,” 2008.
- [100] L. Sokolowski and T. Kozłowski, “Assessment of Two-Phase Critical Flow Models Performance in RELAP5 and TRACE Against Marviken Critical Flow Tests.”
- [101] P. Domitr and M. Włostowski, “The use of machine learning for inverse uncertainty quantification in TRACE code based on Marviken experiment,” *Nucl. Eng. Des.*, vol. 384, no. June, 2021.
- [102] P. Virtanen *et al.*, “{SciPy} 1.0: Fundamental Algorithms for Scientific Computing in Python,” *Nat. Methods*, vol. 17, pp. 261–272, 2020.
- [103] P. Domitr, M. Włostowski, R. Laskowski, and R. Jurkowski, “Comparison of inverse uncertainty quantification methods for critical flow test,” *Energy*, vol. 263, p. 125640, 2023.
- [104] L. E. Hochreiter, “FLECHT SEASET program. Final report,” 1985.
- [105] I. N. Laboratory, “LOFT Systems and Test Description.”
- [106] Los Alamos National Laboratory, “TRAC-M/FORTTRAN 90 (Version 3.0) User’ S Manual,” 2001.
- [107] P. Domitr, L.-Y. Cheng, P. Kohut, and J. Ramsey, “Development of Trace model based on TRAC-M input for PWR loca analysis,” in *17th International Topical Meeting on Nuclear Reactor Thermal Hydraulics, NURETH 2017*, 2017, vol. 2017-Sept.
- [108] P. Domitr, M. Malicki, and L.-Y. Cheng, “Uncertainty assessment of LOCA scenario for trace model based on TRAC-M input of PWR reactor,” in *18th International Topical Meeting on Nuclear Reactor Thermal Hydraulics, NURETH 2019*, 2019.
- [109] R. P. Martin and W. T. Nutt, “Perspectives on the application of order-statistics in best-estimate plus uncertainty nuclear safety analysis,” *Nucl. Eng. Des.*, vol. 241, no. 1, pp. 274–284, 2011.
- [110] J. D. Martin and T. W. Simpson, “Use of Kriging Models to Approximate Deterministic Computer Models,” *AIAA J.*, vol. 43, no. 4, pp. 853–863, 2005.
- [111] A. I. J. Forrester and A. J. Keane, “Recent advances in surrogate-based optimization,” *Prog. Aerosp. Sci.*, vol. 45, no. 1, pp. 50–79, 2009.

9. Aneks

9.1. Metamodele

Metamodele tworzy się na podstawie wykonania niewielkiej liczby symulacji z zastosowaniem kodu obliczeniowego i specjalnie wyselekcjonowanych, reprezentatywnych wartości parametrów. Na tej podstawie algorytm uczenia maszynowego trenuje się do poznania relacji między wartościami wielkości wejściowej oraz wielkości wyjściowej. Metamodela dzięki uproszczeniu zależności opisanej pierwotnie modelem i wymagającej obliczeń z zastosowaniem kodu obliczeniowego, posiadają o wiele mniejszą złożoność obliczeniową i są mniej czasochłonne. Metamodela są więc, wykorzystywane między innymi w metodach wstecznej kwantyfikacji niepewności jako substytutu do wykonywania obliczeń z wykorzystaniem kodu obliczeniowego [27]–[29]. Mogą być też wykorzystywane w obliczeniach wrażliwości czy pracach optymalizacyjnych. Spośród wielu typów metamodeli wykorzystywanych w deterministycznym modelowaniu komputerowym najczęściej stosowane są metamodela procesu Gaussowskiego[32]. Proces Gaussowski to taki proces stochastyczny (czyli zbiór zmiennych losowych indeksowanych w czasie bądź przestrzeni), dla którego każdy skończeniowymiarowy rozkład jest gaussowski. Model procesu gaussowskiego jest modelem liniowej regresji, uwzględniającym korelację pomiędzy modelem regresji a obserwacjami (danymi eksperymentalnymi) przy wyznaczaniu residuów. W ogólnej matematycznej postaci model procesu gaussowskiego przyjmuje postać:

$$y(x) = \mathbf{f}^T(x)\beta + z(x) \quad (9.1)$$

Gdzie pierwsza część równania reprezentuje liniową regresję, w tym funkcja $\mathbf{f} = [f_1, f_2, \dots, f_n]^T$ jest określona przez użytkownika, a wektor $\beta = [\beta_1, \beta_2, \dots, \beta_n]^T$ zawiera współczynniki regresji. Druga składnik równania to stacjonarny proces Gaussowski o zerowej wartości średniej oraz kowariancji określonej następującym równaniem:

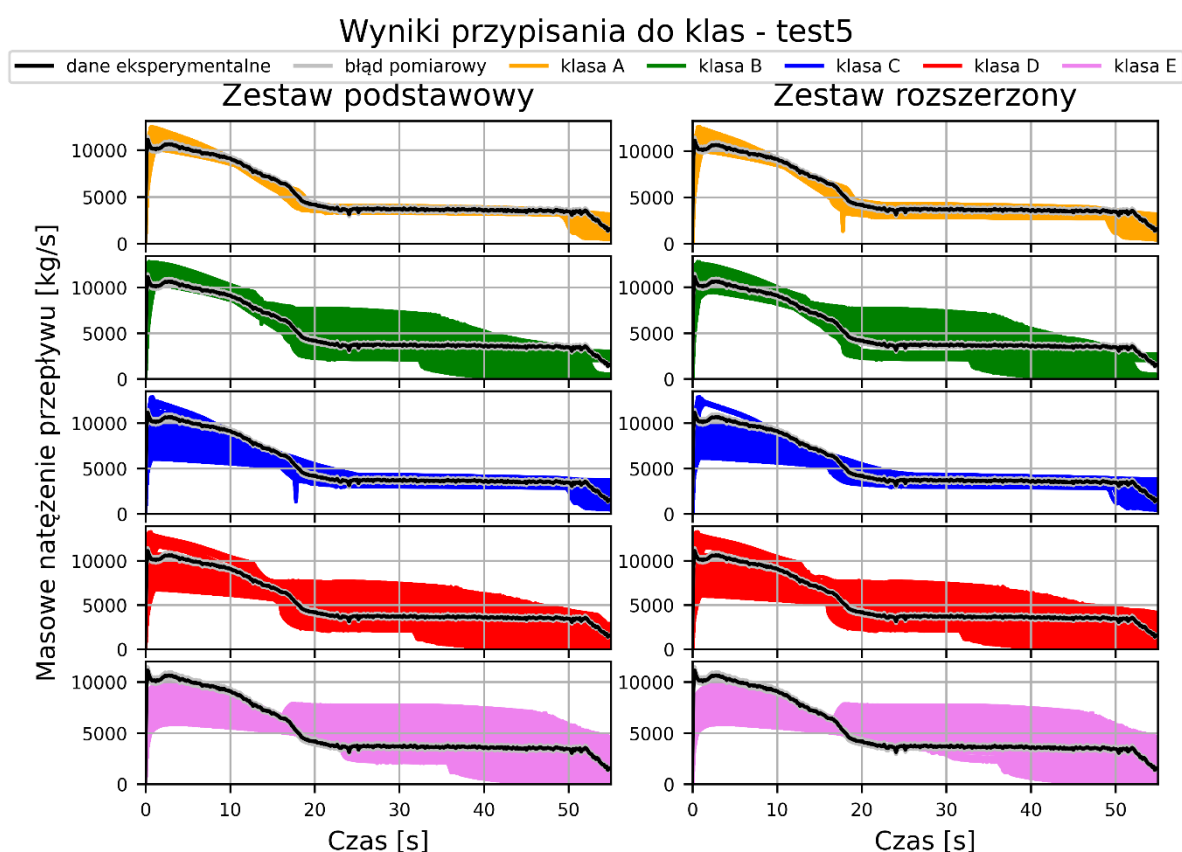
$$\text{Cov}(z(x^i), z(x^j)) = \sigma^2 \cdot K(x^i, x^j) \quad (9.2)$$

Gdzie σ^2 to wariancja a $K(x^i, x^j)$ to funkcja korelacji, zwana również rdzeniem korelacji (ang. “correlation kernel”). Składnik kowariancji w równaniu 9.2 sprawia, że proces gaussowski określony tym równaniem jest modelem procesu stochastycznego.

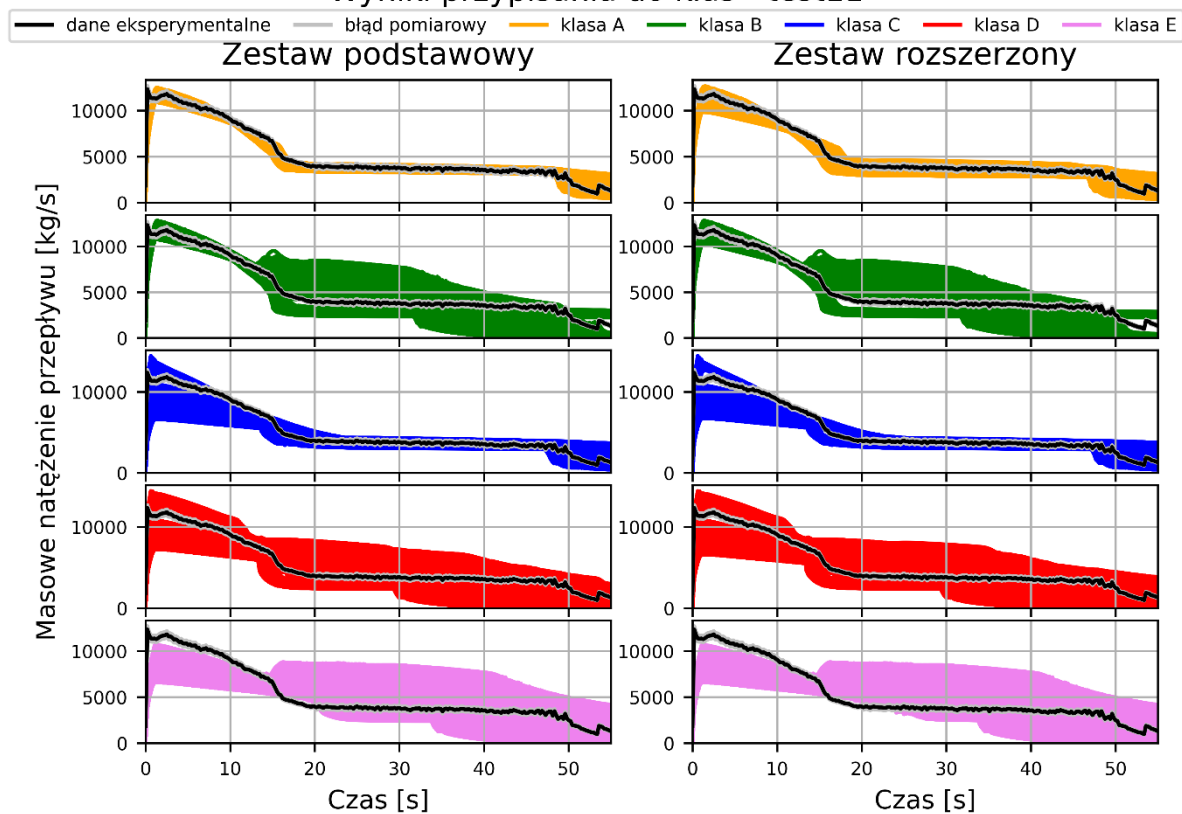
Określenie rdzenia korelacji jest kluczowym elementem budowy procesu Gaussowskiego. Wymagane jest aby rdzeń korelacji był funkcją dodatnio określoną [110], [111]. Jednym z częściej stosowanych podejść jest więc wybór rdzenia korelacji spośród znanych funkcji

korelacji, które są dodatnio określone. Do najpopularniejszych funkcji korelacji zaliczają się liniowe, wykładnicze, Gaussowskie oraz Materna. Spośród tych rdzeni korelacji do zastosowań statystycznych najczęściej wybierane są rdzenie Materna, stąd zaimplementowano je w obecnej metodyce w modelowaniu rozbieżności modelowej. Określenie rdzenia korelacji pozwala na predykcję procesu gaussowskiego dla wybranego x (np. kroku czasowego) określonego równaniem 9.2, jak również na obliczenie błędu średniokwadratowego.

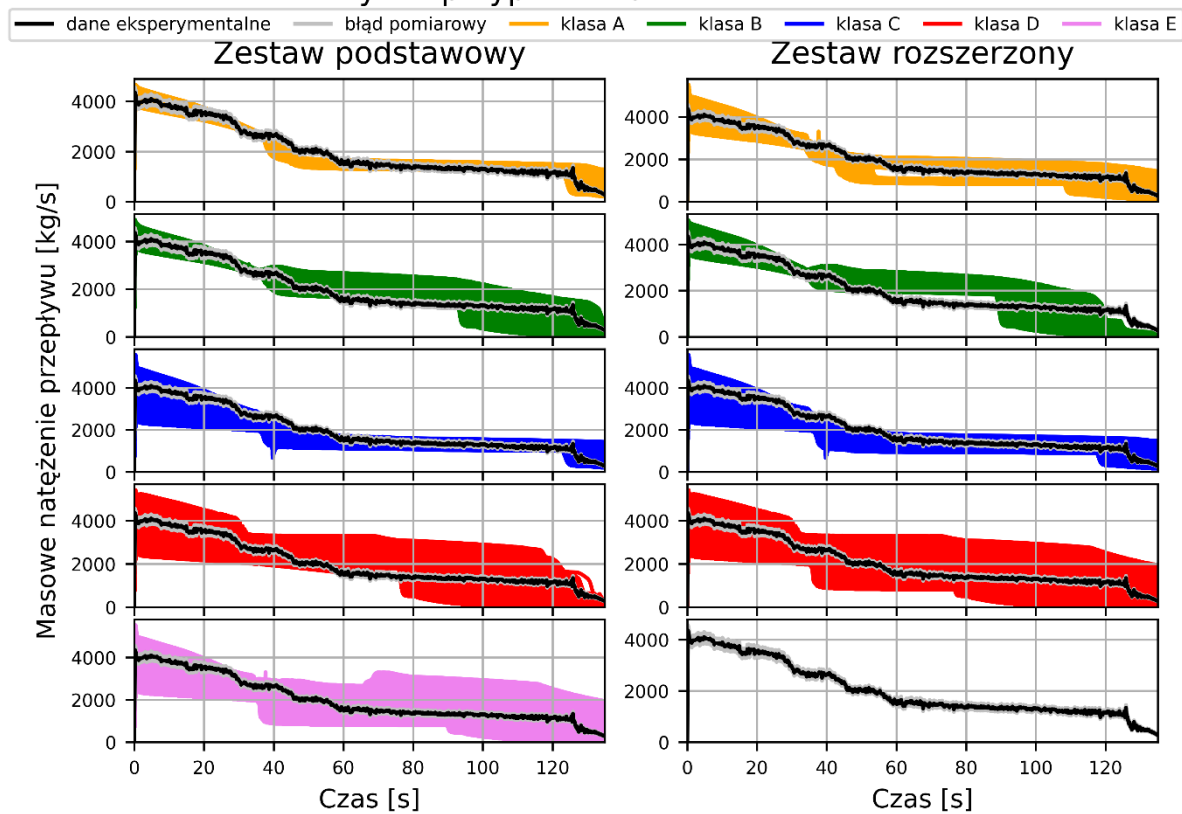
9.2. Zestaw testowy – Marviken – wszystkie klasy dla obu zestawów uczących



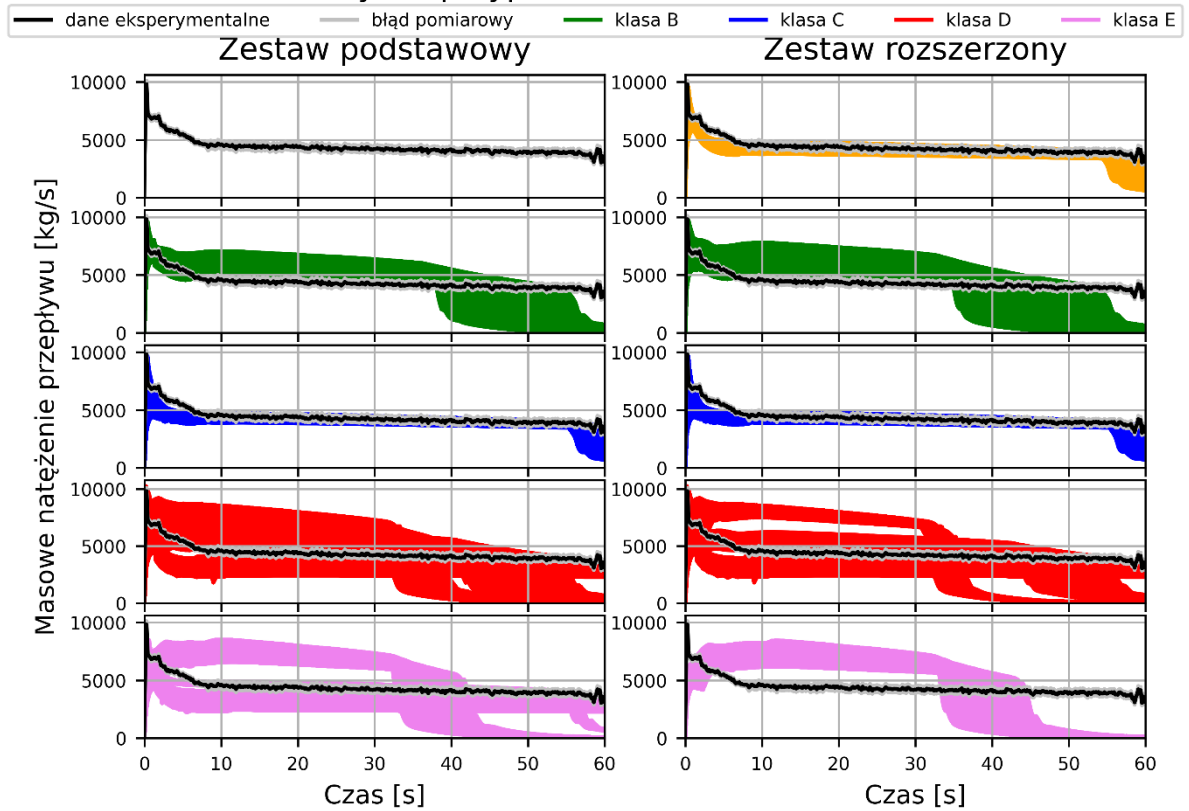
Wyniki przypisania do klas - test11



Wyniki przypisania do klas - test12

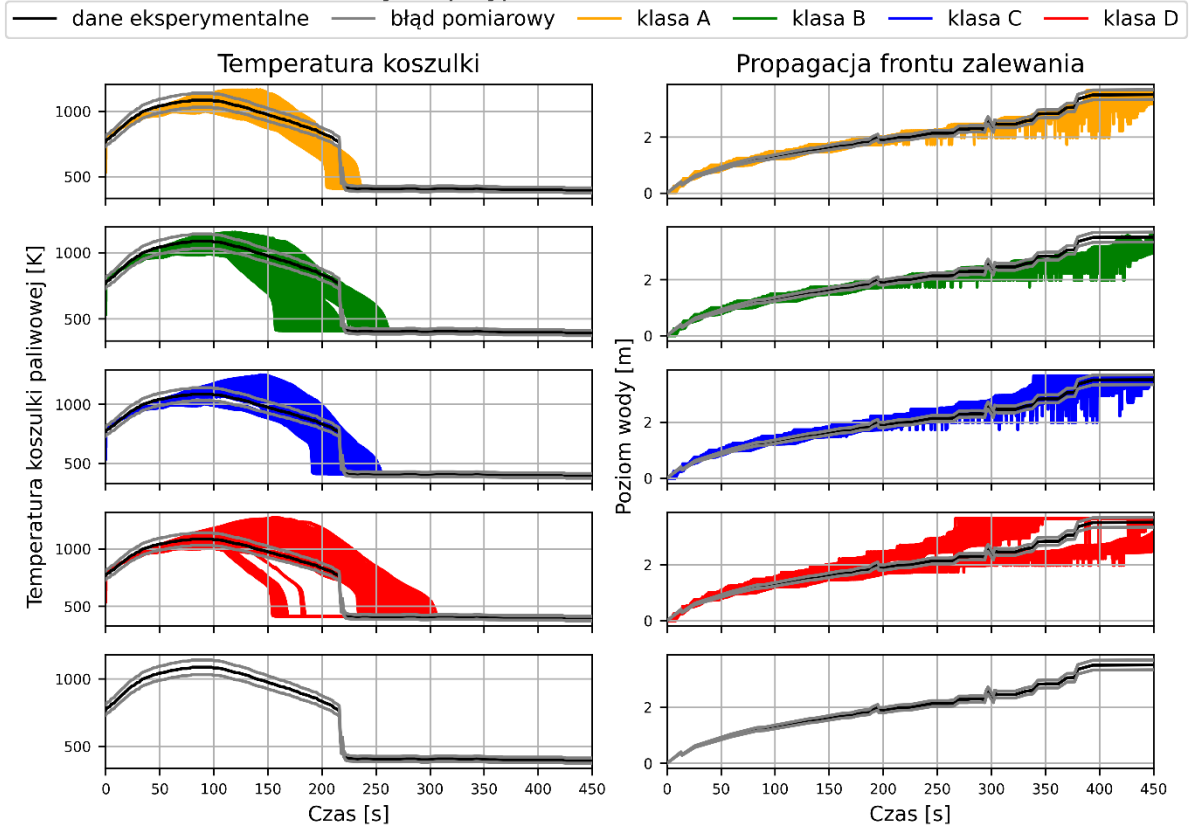


Wyniki przypisania do klas - test20

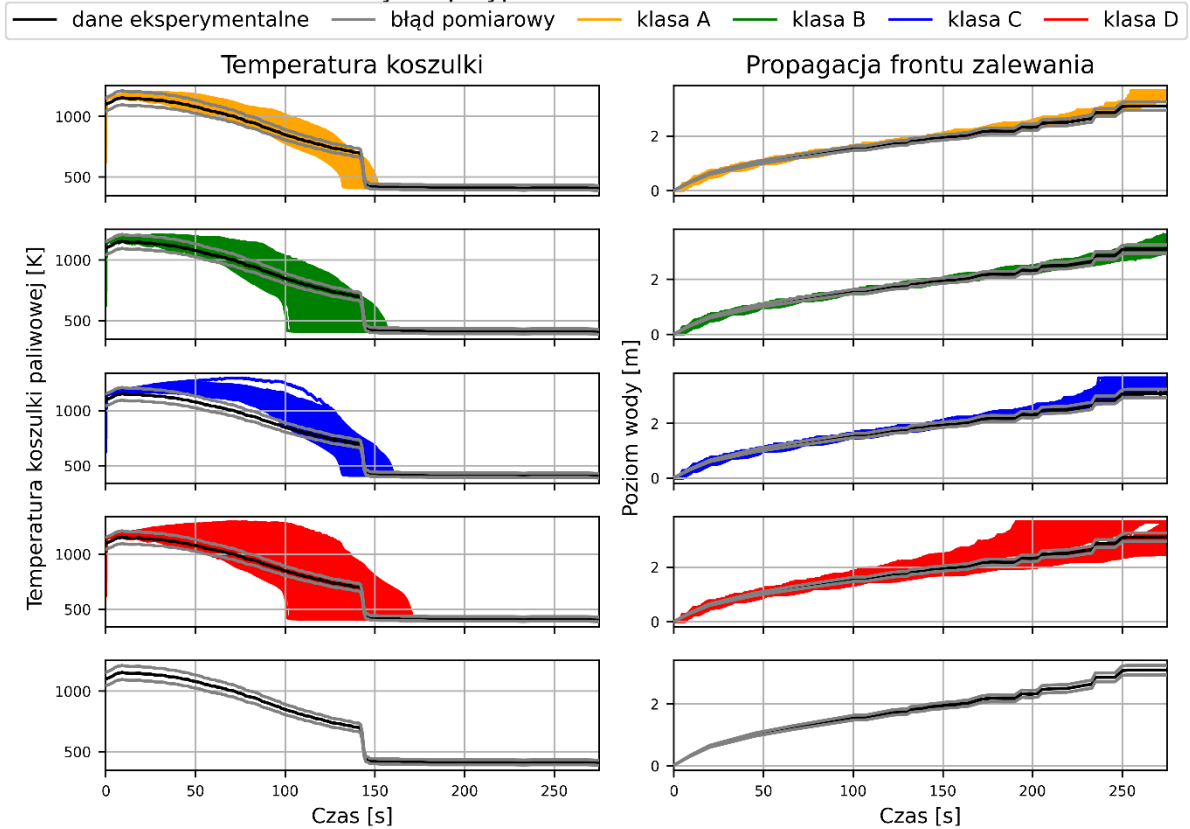


9.3. Zestaw testowy – Flecht Seaset – wszystkie klasy

Wyniki przypisania do klas - test30817



Wyniki przypisania do klas - test31302



Wyniki przypisania do klas - test31504

— dane eksperymentalne — błąd pomiarowy — klasa A — klasa B — klasa C — klasa D — klasa E

