

Wprowadzenie do metodologii opartych na danych w prognozowaniu i zarządzaniu zdrowiem

cze 8, 2025

Prognostyka i zarządzanie zdrowiem (PHM) to dyscyplina inżynierska, której celem jest minimalizacja kosztów konserwacji poprzez ocenę, prognozowanie, diagnozowanie i zarządzanie zdrowiem systemów inżynierskich.

Badania w zakresie prognozowania opartego na danych i zarządzania zdrowiem

Oprócz ogromnego potencjału ekonomicznego w różnych zastosowaniach przemysłowych, PHM ma również dużą wartość badawczą w dziedzinach przetwarzania sygnałów, uczenia maszynowego i eksploracji danych. Badania w zakresie PHM opartego na danych wymagają interdyscyplinarnego zaplecza z zakresu informatyki, przetwarzania sygnałów, statystyki i niezbędnej wiedzy dziedzinowej.

Ogólnie rzecz biorąc, istnieją cztery główne zadania analizy i modelowania danych w PHM:

wstępne przetwarzanie i ekstrakcja cech,
ocena stanu,
diagnostyka,
i prognozowanie

Wstępne przetwarzanie i ekstrakcja cech

Zadanie wstępnego przetwarzania i ekstrakcji cech obejmuje ocenę jakości danych, czyszczenie danych, identyfikację reżimu i segmentację. Chociaż wstępne przetwarzanie nie oferuje bezpośrednio natychmiastowych informacji nadających się do wykorzystania, jest to krytyczny krok i wymaga zarówno wiedzy z danej dziedziny, jak i umiejętności przetwarzania danych, aby zachować cenne części danych, usuwając jednocześnie niepożądane komponenty.

Ocena stanu

Zadanie oceny stanu polega na oszacowaniu i określeniu stanu zdrowia zasobu poprzez analizę zebranych danych. Jeśli istnieją dane o stanie awarii, można wygenerować wartość ufności (CV), aby wskazać prawdopodobieństwo awarii zasobu. Jednak jeśli dane o awarii zasobu nie są dostępne, ocena stanu może zostać przekształcona w problem monitorowania degradacji w przypadku stopniowych usterek lub problem wykrywania usterek w przypadku nagłych usterek.

Diagnostyka

W PHM diagnostyka odnosi się do klasyfikacji różnych trybów awarii poprzez wyodrębnienie sygnatur usterek z danych. Na przykład w przypadku maszyn wirujących proces ten polega na zwiększeniu stosunku sygnału do

szumu sygnałów wibracyjnych i wyodrębnieniu komponentów cyklo-stacjonarnych, które mogą reprezentować wady w niektórych komponentach maszyny. Zbiór różnych funkcji można wykorzystać wraz z algorytmem klastrowania lub klasyfikacji w celu opracowania modelu opartego na danych do diagnostyki usterek maszyn.

Prognostyka

Zadaniem prognozowania jest przewidywanie stanu zdrowia aktywów. Jeśli pożądana jest krótkoterminowa prognoza, modelowanie szeregów czasowych jest często wykorzystywane do przewidywania, kiedy maszyna przekroczy próg. Jeśli preferowana jest długoterminowa prognoza, problem staje się prognozą pozostałego okresu użytkowania (RUL) z wieloma dostępnymi narzędziami uczenia maszynowego i statystyki. Zakres ufności będzie musiał zostać zdefiniowany dla takich prognoz, ponieważ wydajność maszyny jest również w dużym stopniu zależna od wzorca użytkowania i odpowiednich działań konserwacyjnych, które zostaną podjęte.

Metodologia

Inteligentne przetwarzanie koncentruje się na wykorzystaniu narzędzi konwersji danych na informacje w celu przekształcenia surowych danych aktywów w informacje nadające się do wykorzystania, takie jak wskaźniki stanu, zalecenia konserwacyjne i prognozy wydajności. Wszystkie te informacje są kluczowe dla użytkowników, aby w pełni zrozumieć bieżącą sytuację monitorowanego zasobu i podejmować zoptymalizowane decyzje.

Dostępne narzędzia przetwarzania danych na informacje do inteligentnego przetwarzania obejmują: modele oparte na fizyce, modele statystyczne i algorytmy uczenia maszynowego/eksploracji danych. Spośród wszystkich trzech opcji uczenie maszynowe/eksploracja danych ma swoje korzenie w praktykach komputerowych i ma wiele zalet w zastosowaniach przemysłowych [4–6]:

(1) Mniejsze wymagania dotyczące wiedzy domenowej [7]: bez budowania dokładnego modelu matematycznego dla systemu fizycznego, uczenie maszynowe może również wyodrębnić przydatne informacje, obserwując pary danych wejściowych i wyjściowych, tak jak robią to modele fizyczne. Ta cecha sprawia, że uczenie maszynowe jest przydatne w przypadku złożonych systemów inżynierskich i procesów przemysłowych, w których wcześniejsza wiedza jest niewystarczająca do zbudowania zadowalających modeli fizycznych.

(2) Skalowalność dla różnych zastosowań [8].

(3) Łatwa implementacja: w porównaniu z modelami opartymi na fizyce, uczenie maszynowe jest bardziej odpowiednie do obsługi dużych zestawów danych, ponieważ wymaga mniej zasobów obliczeniowych.

Algorytmy

Algorytmy PHM odnoszą się do modeli opartych na danych, które służą jako rdzeń obliczeniowy do przekształcania cech w znaczące informacje. Na rys. 3 opisano przykład tego procesu oceny stanu silników indukcyjnych. Proces będzie podobny w przypadku diagnostyki i prognozowania, ale ostateczne narzędzie wizualizacji będzie się różnić w zależności od celu użytkowników.

Tabela 1 Podsumowanie metod wstępnego przetwarzania danych

Metoda Zalety Wady

1 Kontrola jakości danych

Wymaga wcześniejszej znajomości typu sygnału, ale jest skuteczna w szczególności w przypadku walidacji sygnału wibracji

Progi są potrzebne do określenia, czy uwzględnić sygnał w analizie

2 Identyfikacja reżimu

Identyfikacja reżimu jest ważna dla opracowania bazowych zestawów danych w każdych warunkach pracy

Algorytmy wstępnego przetwarzania danych

Podsumowanie narzędzi wstępnego przetwarzania danych wraz z zaletami i wadami każdej metody przedstawiono w Tabeli 1. W wielu zastosowaniach jedna lub więcej z tych metod jest potrzebnych do zapewnienia, że dane nadają się do dalszego przetwarzania i opracowywania algorytmów. Kontrola jakości danych jest szczególnie ważna, aby upewnić się, że nie występują błędy czujnika lub akwizycji danych. Jednak opracowanie skutecznego algorytmu wymaga pewnej wcześniejszej wiedzy na temat charakterystyki sygnału lub dystrybucji sygnału.

W przypadku bardziej skomplikowanych reżimów pracy może być konieczne posiadanie bardziej zaawansowanych technik identyfikacji warunków pracy. Informacje o reżimie pozwalają na opracowanie bazowych zestawów danych w każdym stanie pracy i przeprowadzenie uczciwego porównania oraz lokalnego modelu stanu w każdym stanie pracy.

Usuwanie wartości odstających z mierzonych danych jest również kluczowe w niektórych zastosowaniach, ponieważ obecność wartości odstających może znacząco wpłynąć na wynik analizy. Istnieją różne metody usuwania wystąpień wartości odstających z sygnału lub wyodrębnionych cech. Jednakże metody te opierają się wyłącznie na dystrybucji danych lub cechach, a doświadczenie inżynierskie i wiedza domenowa mogą być wykorzystane do dalszego udoskonalenia algorytmu usuwania wartości odstających.

Algorytmy ekstrakcji cech

Dostępnych jest wiele metod i algorytmów do ekstrakcji cech lub cech z mierzonych sygnałów, a przegląd niektórych dostępnych technik przedstawiono w Tabeli 2.

W przypadku sygnałów o wysokiej częstotliwości, takich jak drgania lub prąd, istnieją dobrze ugruntowane metody przetwarzania sygnałów i ekstrakcji cech w celu ekstrakcji informacji z reprezentacji sygnału w domenie czasu i częstotliwości [10].

W przypadku łożysk tocznych, wałów mechanicznych i kół zębatych istnieje kilka konkretnych metod przetwarzania w celu ekstrakcji cech degradacji dla tych komponentów [11]. Chociaż korzystne jest stosowanie metod ekstrakcji cech specyficznych dla komponentów dla sygnałów drgań i prądu o wysokiej częstotliwości, wymagają one wyższej częstotliwości próbkowania, większej liczby obliczeń i droższych systemów akwizycji danych.

W przypadku zastosowań, w których monitorowany zestaw sygnałów składa się z sygnałów śladowych, takich jak temperatura, ciśnienie lub inne sygnały sterujące, zalecany byłby inny zestaw algorytmów ekstrakcji cech. Algorytmy przetwarzania oparte na resztkach, takie jak sieci neuronowe autoasocjacyjne lub metody oparte na głównych składnikach, to przykłady metod, których można użyć do przetwarzania sygnałów śledzenia [12]. W przypadku tego typu algorytmów, linia bazowa jest ustalana na podstawie normalnej pracy maszyny.

Ta linia bazowa jest używana do porównywania przewidywanych wartości czujników i przetwarzania reszt jako znaku dryfu w wydajności maszyny. Ekstrakcja różnych parametrów statystycznych jest prostym, ale skutecznym podejściem do charakteryzowania stanu systemu z dostępnych sygnałów sterownika.

W wielu przypadkach można uzyskać więcej informacji, ekstrahując statystyki w różnych przedziałach czasowych, a przedział czasowy może reprezentować inny ruch lub działanie wykonywane przez monitorowany system [13]. Jeśli informacje kontekstowe dotyczące sygnałów procesu nie są dostępne, odpowiednią alternatywą jest ekstrakcja statystyk czasowych bez żadnej segmentacji.

Algorytmy oceny stanu i wykrywania anomalii

Lista najczęściej stosowanych algorytmów oceny stanu maszyny znajduje się w Tabeli 3. Najprostszym podejściem do oceny stanu jest wyodrębnienie metryki stanu na podstawie ważonego sumowania wartości cech. Ta metryka stanu jest prosta do obliczenia, a następnie można wyprowadzić progi statystyczne wykrywania degradacji na podstawie rozkładu wartości stanu [14].

Do określenia stanu zdrowia monitorowanego systemu lub komponentu można użyć różnych odległości od normalnych metryk stanu. Odległość Mahalanobisa i analiza głównych składowych oparta na statystykach T2 Hotellinga to metryki odległości, które uwzględniają relację kowariancji między zmiennymi; jednak powszechnie stosuje się również euklidesowe i inne metryki odległości [15]. Metody oparte na odległości nie uwzględniają, czy cechy są wyższe czy niższe od normalnych.

Ma to potencjalne wady, ponieważ system może mieć niższe drgania niż normalne, a to nadal powodowałoby wyższą wartość stanu zdrowia opartą na odległości i stan anomalii.

Prostym, ale skutecznym podejściem do wykrywania anomalii jest wykorzystanie statystycznego testowania hipotez. Sekwencyjny test ilorazu prawdopodobieństwa, test permutacji rang i test T to przykłady niektórych z powszechniej stosowanych testów hipotez do wykrywania anomalii [16]. Inne metody oparte na wykrywaniu anomalii obejmują klasyfikator jednoklasowy, taki jak algorytm opisu danych wektora nośnego (SVDD) [17].

W odniesieniu do SVDD jedną z wad jest brak wskazówek w literaturze na temat tego, jakich funkcji jądra lub ustawień używać w danej aplikacji. Metody oparte na regresji lub sieci neuronowej są szczególnie skuteczne, jeśli dostępne są wystarczające dane do opracowania modeli regresji. Sieć neuronowa lub model regresji można wykorzystać do zapewnienia mapowania między wartościami cech a wartością stanu zdrowia lub rozmiarem defektu [18]. Zdolność modelu do uogólniania zwykle wymaga wielu zestawów danych treningowych.

Algorytmy diagnostyki stanu zdrowia

Do analizy przyczyn źródłowych i diagnostyki istnieje wiele różnych metod i algorytmów do wykonania tego zadania, a próbka niektórych z najczęściej stosowanych metod została wymieniona w Tabeli 4. Włączenie wiedzy inżynierskiej i doświadczenia do algorytmu diagnostycznego sprawia, że wykorzystanie rozmytych funkcji przynależności i reguł jest atrakcyjną techniką [19]. Jednak staje się trudniejsze wykorzystanie rozmytego algorytmu diagnostycznego w przypadku nowych aplikacji, w których nie ma wystarczającego doświadczenia w zakresie trybów awarii i ich sygnatur.

Użycie algorytmu klasyfikacji jest popularną alternatywą, jeśli istnieją dane z wielu stanów zdrowia, w tym stanu bazowego i kilku różnych trybów awarii, które mogą wystąpić.

Użycie sieci neuronowych, maszyn wektorów nośnych i algorytmu Naïve Bayes to jedne z najczęściej stosowanych algorytmów klasyfikacji do monitorowania stanu maszyn [20]. Poprzez poznanie relacji między wyodrębnionymi cechami a sygnaturami bazowymi i awarii, metoda klasyfikacji może dokładnie zdiagnozować i oznaczyć stan zdrowia monitorowanego systemu.

Inna metoda diagnostyki obejmuje wykorzystanie Bayesowskiej Sieci Przekonań (BBN), która zapewnia sieć reprezentującą związek przyczynowy między mierzoną zmienną a różnymi trybami awarii lub warunkami systemu, które mogą wystąpić [21].

Algorytmy prognostyczne

Przykład najczęściej stosowanych algorytmów przewidywania pozostałego okresu użytkowania przedstawiono w Tabeli 5, wraz z zaletami i wadami każdej metody.

Metody oparte na dopasowaniu krzywych są stosunkowo proste w zastosowaniu, ponieważ nie wymagają znacznej ilości danych treningowych ani szczegółowego modelu fizycznego opisującego postęp usterki.

Metody oparte na sieci neuronowej lub regresji mogą bezpośrednio powiązać wartości cech z pozostałym okresem użytkowania monitorowanego komponentu lub systemu [22]. Metody te wymagają znacznych danych treningowych do nauczenia się tej relacji, a uzyskanie wielu zestawów danych od uruchomienia do awarii nie jest wykonalne w wielu zastosowaniach.

Algorytm prognostyczny oparty na podobieństwie to unikalna metoda, która dopasowuje poprzednie wzorce degradacji do bieżącego wzorca degradacji monitorowanego systemu [23]. Algorytm prognostyczny oparty na podobieństwie może być dość dokładny; jednakże wymaga kilku przebiegów do zestawów danych awarii w celu uzyskania biblioteki do wykonania dopasowania trajektorii degradacji. W przeciwieństwie do wcześniej opisanych algorytmów prognostycznych opartych na danych, włączenie fizyki awarii do algorytmu filtrowania stochastycznego jest skutecznym podejściem.

niezależnie od tego, czy równania dynamiczne propagacji awarii są liniowe czy nieliniowe, a także czy szum pomiarowy jest gaussowski czy nie, można użyć tej grupy metod [24]. Stosowanie algorytmów predykcyjnych opartych na modelach z wykorzystaniem filtrowania stochastycznego wiąże się z pewnymi potencjalnymi wyzwaniem. Tylko dla podzbioru aplikacji istnieją ustalone modele opisujące mechanizm awarii.

W poniższej sekcji przedstawiono studium przypadku, które bardziej szczegółowo wyjaśnia, w jaki sposób kompleksowy system PHM można zastosować w rzeczywistych zastosowaniach i w jaki sposób przemysł może skorzystać z takiego systemu.

