



DeePMO: An iterative deep learning framework for high-dimensional kinetic parameter optimization[☆]

Pengxiao Lin^{a,b}, Yuntian Zhou^c, Zhiwei Wang^{a,b}, Weizong Wang^d, Zheng Chen^c,
Zhi-Qin John Xu^{a,b},^{*}, Tianhan Zhang^{d,e},^{**}

^a Institute of Natural Sciences, MOE-LSC, Shanghai Jiao Tong University, Shanghai, 200240, China

^b School of Mathematical Sciences, Shanghai Jiao Tong University, Shanghai, 200240, China

^c HEDPS, SKLTCS, College of Engineering, Peking University, Beijing, 100871, China

^d School of Astronautics, Beihang University, Beijing, 100191, China

^e Key Laboratory of Spacecraft Design Optimization and Dynamic Simulation Technologies, Ministry of Education, Beijing, 102206, China

ARTICLE INFO

Keywords:

Deep neural network
Chemical kinetics
Parameter optimization
Machine learning
Iterative strategy

ABSTRACT

This study introduces Deep learning-based kinetic model optimization (DeePMO), a novel approach for optimizing parameters in chemical kinetic models. The primary challenge lies in mapping high-dimensional kinetic parameters to comprehensive performance metrics derived from diverse numerical simulations, including ignition delay time, laminar flame speed, heat release rate, and temperature-residence time distributions in perfectly stirred reactors. We propose an iterative sampling-learning-inference strategy to efficiently explore high-dimensional parameter spaces. The approach features a hybrid deep neural network (DNN) architecture that combines a fully connected network for non-sequential data with a multi-grade network for sequential data, enabling effective utilization of performance metrics with varying distribution characteristics. DeePMO's effectiveness and versatility was validated across multiple fuel models, including methane, ethane, butane, n-heptane, n-pentanol, ammonia, ammonia/hydrogen, and their mixtures, with parameter counts ranging from tens to hundreds. The validation demonstrated successful optimization in all test cases and confirmed the method's flexibility in incorporating both direct experimental measurements and simulated data from benchmark chemistry models. An ablation study highlighted the critical role of DNN in guiding data sampling and optimization processes, while additional comparative experiments examined hyperparameter effects. This work provides a valuable tool for kinetic parameter optimization and offers insights for applying machine learning algorithms in combustion research.

1. Introduction

Combustion chemical kinetics is a complex nonlinear system involving many elementary reactions and intermediate species, where calibrating rate parameters is essential [1]. For detailed models, the individual rate measurement and reaction-rate theory are born with uncertainty, let alone estimation for analogous reactions. The limited understanding of the physical processes can further complicate the model formulation [2,3]. Thus, parameter optimization becomes necessary for uncertainty quantification and calibrating with experimental data [4]. The optimization technique becomes even more important for developing reduced models, as extra truncation errors are involved when removing redundant species and reactions from the full model.

The modified reaction pathway structure requires fine-tuned reaction rate parameters to better predict quantities of interest.

Approaches to rate parameter tuning include genetic algorithms (GA) [5,6], as demonstrated in the optimization of methane flames [7], and hydrogen/nitrogen/oxygen combustion [8], where GA achieved accurate predictions of flame speeds, species profiles, and ignition delays while streamlining mechanisms for computational efficiency. GA is a metaheuristic optimization method that employs principles of natural evolution to search for optimal parameter sets by iteratively evaluating and evolving populations of candidate solutions. Sensitivity analysis serves as a diagnostic tool to identify influential parameters by perturbing them and assessing changes in model response. Current rate optimization methods often build upon sensitivity

[☆] This article is part of a Special issue entitled: 'AI for Combustion' published in Applications in Energy and Combustion Science.

^{*} Corresponding author at: Institute of Natural Sciences, MOE-LSC, Shanghai Jiao Tong University, Shanghai, 200240, China.

^{**} Corresponding author at: School of Astronautics, Beihang University, Beijing, 100191, China.

E-mail addresses: xuzhiqin@sjtu.edu.cn (Z.-Q.J. Xu), thzhang@buaa.edu.cn (T. Zhang).

analysis, extending it to higher orders to solve the inverse problem of uncertainty quantification. A forward problem is constructed first to quantify the impact of uncertainties in kinetic parameters on the model's performance, such as the accuracy of predicting ignition delay time, laminar flame speed, or extinction stretch rate. The quantitative correlation between the rate parameters and the model predictions is formulated as the solution mapping and response surface. Advanced mathematical tools from informatics and data science are introduced to represent the high-dimensional nonlinear mapping functions, including the Bayesian approach [9], polynomial chaos expansion (PCE) [4], and high-dimensional model representation (HDMR) [10,11]. Concurrently, popular tools such as *OptiSMOKE++* [12] and *Optima++* [13,14] have been developed to facilitate the optimization of kinetic mechanisms against experimental data. *OptiSMOKE++* leverages the DAKOTA toolkit to employ derivative-free optimization methods, such as DIRECT, MADS, Solis-Wets, and pattern search, while bounding parameters within their uncertainties. *Optima++* uses FOCTOPUS optimization algorithm [15] combined with principal component analysis (PCA) to minimize a least squares error function.

Powerful mathematical tools facilitate the kinetic parameter optimization methods. In recent years, deep learning has emerged as a highly successful mathematical tool in various scientific and engineering domains [16–22]. The explosive development of deep learning can be attributed to three main factors [23]: the availability of massive datasets, advancements in GPU-based computation power, and algorithmic improvements. These factors benefit combustion research as well. Deep learning is particularly suited for studies related to reaction kinetics due to its ability to uncover intricate structures within reaction networks and identify high-order correlations in high-dimensional data. When it comes to rate optimization, employing a deep neural network as a surrogate model using rate parameters as inputs and quantities of interest as outputs is a good demonstration in [24,25]. This training process can be seen as solving the forward problem, while backpropagation serves as the inverse problem. The pioneering works by [26] showed the advantage of using neural networks as surrogate models to replace PCE or HDMR, while the one-shot sampling strategy and one-hidden-layer structure might limit the neural network approximation ability of high-dimensional mapping functions. In [27], researchers implemented neural ODE into kinetic optimization and demonstrated its strong performance using temporal experiment data. Owoyele and Pal [28] introduced ChemNODE, embedding neural ODEs into chemical solvers to ensure trajectory fidelity during hydrogen-air autoignition predictions and maintain gradient propagation for optimization. Fedorov et al. [29] proposed Kinetics-Constrained Neural ODEs, embedding thermodynamic knowledge into ANN architectures to enable reliable kinetic modeling with limited experimental data. Zhang et al. [30] advanced adaptive ANN training for high-dimensional uncertainty minimization in FFCM-2 mechanisms, outperforming polynomial methods. Recently, Wang et al. [31] and Chen et al. [32] replaced MCMC with variational inference guided by ANN surrogates, achieving speedup in methanol kinetic optimization while resolving parameter covariance constraints. Zhou et al. [33] proposed the OptEx framework to optimize combustion models through experimental design and data clustering. Li et al. [34] developed the CODENN algorithm, integrating neural ODEs and Cantera to optimize the high-fidelity battery venting gas mechanism. However, developing a general and effective method of optimization of kinetic parameters presents significant challenges that remain largely unexplored. These challenges stem from the need to accommodate diverse performance metrics within a unified framework while navigating high-dimensional parameter spaces. An ideal solution must be robust across different chemical systems, scalable to handle varying parameter quantities, compatible with heterogeneous data types, and capable of maintaining physical interpretability of the optimized parameters.

This study introduces DeePMO (Deep learning-based kinetic model optimization), a novel framework designed to overcome critical limitations in chemical kinetic model development. DeePMO tackles the

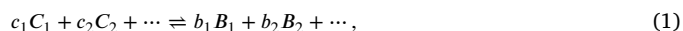
curse of dimensionality by enabling the simultaneous optimization of hundreds of kinetic parameters—a significant leap from conventional methods typically limited to dozens. Its core innovation is an iterative sampling strategy that trains sequential local DNN surrogates for improved sample efficiency and accuracy, paired with a hybrid DNN architecture integrating fully connected networks for non-sequential QoIs (e.g., ignition delay time) and multi-head networks for sequential data (e.g., temperature profiles in perfectly stirred reactors). This unified approach accommodates diverse QoIs, effectively utilizing both high-fidelity simulations and sparse experimental measurements, with rigorous validation on independent datasets to ensure reproducibility, robustness, and generalization.

In the subsequent sections, we substantiate these claims. The methodology section details the DeePMO methodology, including the hybrid DNN architecture for diverse QoIs, iterative local surrogate training for sample efficiency, and rigorous uncertainty constraints to ensure physical plausibility, with explicit norms and error metrics for reproducibility. The results and discussion section validates DeePMO across various fuel models, including n-heptane, ammonia, and others, demonstrating its robust performance. Through a detailed ablation analysis, we prove the critical role of the deep neural network in achieving these capabilities and provide an in-depth kinetic analysis of the optimization results. Finally, we conclude by summarizing the key contributions of this work.

2. Methodology

2.1. Problem formulation and deep neural network design

A general chemical reaction can be written as



where c_i 's and b_i 's are stoichiometric coefficients and C_i 's and B_i 's are reactants and products, respectively.

The rate coefficient (or rate constant) of a chemical reaction can be expressed in the form of the Arrhenius equation, i.e.,

$$k = AT^b \exp\left(-\frac{E_a}{k_B T}\right), \quad (2)$$

where A , b , E_a are the pre-exponential factor, the temperature dependency exponent, and the Arrhenius energy of activation and k_B is denoted as the Boltzmann constant. These kinetic parameters are crucial for models to predict QoIs accurately in numerical simulations, while the comprehensive relationship between kinetic parameters and overall model performance is challenging to depict.

As illustrated in Fig. 1(a), the optimization process considers five distinct QoIs for evaluating the chemical kinetic model, which are defined as follows:

Ignition Delay Time (IDT): Defined as the time to the intersection of the baseline and the tangent at the maximum slope of the OH concentration profile in an adiabatic, constant-volume reactor simulating shock tube experiments. Due to spanning orders of magnitude, IDTs are analyzed on a logarithmic scale, with errors computed as absolute logarithmic differences (e.g., $|\log(\tau_{\text{sim}}/\tau_{\text{exp}})|$).

Laminar Flame Speed (LFS): The propagation speed of a one-dimensional, planar flame through a stationary premixture.

PSRT and PSRex: Steady-state temperature as a function of residence time (PSRT) and the residence time at flame extinction (PSRex) in a perfectly stirred reactor (PSR).

Mole fraction in PSR: Species mole fractions in a PSR at fixed residence time.

As depicted in Fig. 1(b), we employ a DNN to model the response surface between pre-exponential factors A and QoIs. While the method can be extended to model relationships with other kinetic parameters, such as activation energy, this study focuses specifically on optimizing pre-exponential factors to demonstrate the methodology's effectiveness.

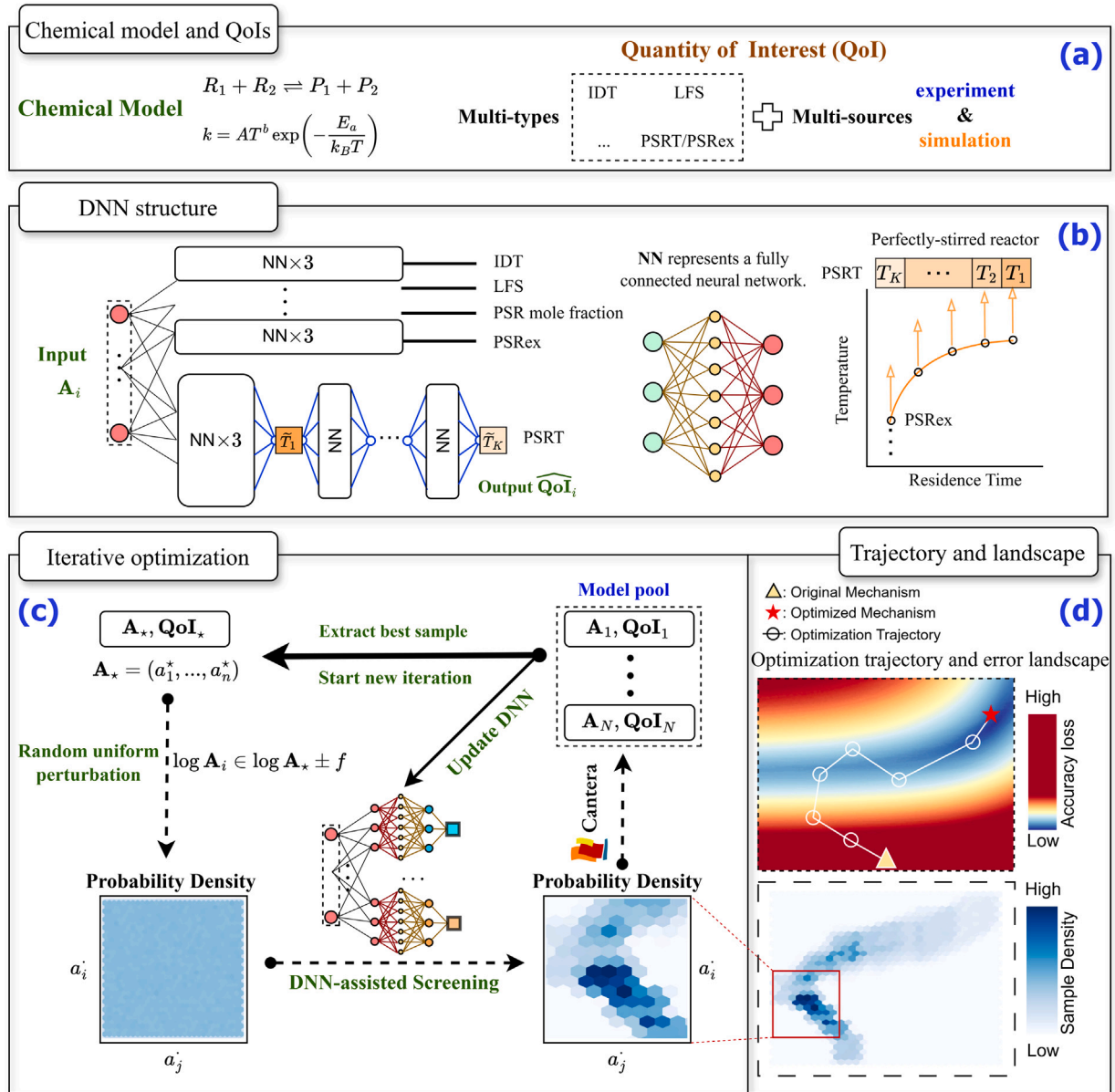


Fig. 1. Overview of the DeePMO methodology through a four-part flow chart. Fig. (a) Benchmark data acquisition: DeePMO obtains optimization benchmarks from either experimental data or detailed model simulations. Fig. (b) Hybrid DNN architecture: fully connected network for non-sequential data (IDT, LFS, HRR, and PSRex) and multi-grade networks for sequential data (PSRT). Every layer of DNN is a fully connected neural network with ReLU activation. Fig. (c) Iterative sampling-training-optimization process. In every iteration, we generate the model pool based on the best mechanism of the previous iteration and the assistance of DNN. Fig. (d) shows the visualization of the optimization process across iterations, presented in two complementary subfigures. The coordinate system represents two representative pre-exponential factors from the kinetic model. The upper subfigure demonstrates the progressive improvement in DNN accuracy throughout the iterations, culminating in identifying the global minimum for model errors. The optimization trajectory is mapped as a continuous path through the parameter space. The lower subfigure illustrates the spatial distribution of sampled data points along this optimization trajectory, revealing the evolution of the sampling strategy as the algorithm converges toward optimal parameters.

The architecture implements distinct subnetworks for different types of QoIs. A standard fully connected network structure suffices for non-sequential data such as IDT. However, PSRT's sequential nature necessitates a specialized network design. We implement a K -layer sub-network for PSRT calculations, where each layer operates sequentially: the k th layer predicts the PSRT value at time step k and feeds its output to the $(k+1)$ th layer, enabling adaptive prediction refinement through the temporal sequence. All experiments in this study use $K = 20$ time steps.

We distinguish three versions of QoIs in our methodology: (1) $Q(A)$: QoIs calculated through numerical simulation as a function of pre-exponential factors A ; (2) $DNN(A)$: QoIs predicted by the DNN

surrogate model; (3) Q^b : Ground-truth benchmark QoIs, and here represents for the simulation data of detailed mechanism or experiment data. The DNN predicts model performance, with its loss function $\tilde{L}_Q(A)$ defined as the mean square error between $DNN(A)$ and $Q(A)$, where θ denotes the parameters in neural network and λ is treated as the L^2 regularization coefficient in DNN training.

$$\tilde{L}_Q(A) = \|DNN(A) - Q(A)\|_2^2 + \lambda \|\theta\|_2^2 \quad (3)$$

Once converged, the DNN can efficiently predict QoIs for arbitrary combinations of kinetic parameters without requiring computationally expensive numerical simulations to obtain $Q(A)$. The ultimate objective is to identify optimal parameters that minimize the discrepancy

between the optimized model's outcome $Q(A)$ and benchmark values Q^b . This optimization is achieved by minimizing the error function defined in Eq. (4), where a hyperparameter weight vector ω balances the relative importance of different QoIs.

$$\mathcal{L}(A) = \sum_Q \omega_Q \left\| \frac{Q(A) - Q^b}{\sigma} \right\|_{\infty}. \quad (4)$$

We utilize the infinity norm to ensure the optimized mechanism performs uniformly well across all conditions of interest, with each target inversely weighted by its uncertainty σ .

Weight hyperparameters play a pivotal role in the optimization process. Traditional methods often weight QoIs based on experimental uncertainties, assigning lower weights to higher-uncertainty targets. However, this approach has limitations: it lacks flexibility for user-defined priorities, e.g., emphasizing IDT fidelity over LFS or vice versa, and does not directly apply to simulation data, which lack measurement noise but can incorporate propagated uncertainties via condition variability as standard practice. To address these, DeePMO employs tunable weights ω_Q between QoIs to guide optimization, alongside preprocessing for standardization – such as logarithmic transforms for IDTs and normalization by uncertainties σ – ensuring comparable errors across heterogeneous data types. For simulation data, we propagate uncertainties where applicable or set σ to the minimum experimental uncertainty for consistency, with all values documented in the Supplementary Information for reproducibility.

2.2. Iterative sampling-training-inference process

As illustrated in Fig. 1(c), we employ a neural network-assisted iterative process to search for optimal pre-exponential factors. At iteration t , we define several key variables. Let A_i represent the i th model among N optimized models, and the set $\{A_i\}_{i=1}^N$ is denoted as the model pool. Each model contains nr reactions, with individual pre-exponential factors denoted as $a_{i,j}$. Thus, $A_i = \{a_{i,j}\}_{j=1 \dots nr}$ represents the complete pre-exponential factor vector for a single model. These varying kinetic parameters across models enable the DNN to learn the mapping between parameters and model performance. The A^* represents the best-performing model in the current iteration:

$$A^* = \arg \min_i \{\mathcal{L}(A_i) : i = 1, \dots, N\}. \quad (5)$$

The adjustment range for each kinetic parameter is carefully constrained to ensure physical plausibility, guided by uncertainty quantification (UQ). For the j th reaction, we define a maximum allowable adjustment bound, denoted as $f_{bound,j}$, for the logarithm of its pre-exponential factor, a_j . This bound establishes a hard constraint on the parameter search space, such that the adjusted value, a'_j , must remain within $B = [\log a_j^0 - f_{bound,j}, \log a_j^0 + f_{bound,j}]$. The value of this non-tunable bound, $f_{bound,j}$, is determined based on the nature of the kinetic mechanism: (1) For detailed mechanisms with provided UQ, $f_{bound,j}$ is set directly to the uncertainty factor reported for that reaction. (2) For detailed mechanism reactions without a specified UQ, we adopt a conventional approach based on prior studies [27,30] and set $f_{bound,j} = 1.3$. (3) For reduced mechanisms, where original UQ values are no longer applicable, we set a conservative, uniform bound of $f_{bound,j} = 1.5$.

Within these hard boundaries, we introduce a separate hyperparameter, the iterative sampling range f_j to control the perturbation magnitude during the optimization process. Before screening we uniformly sample $[\log a_j^* - f_j, \log a_j^* + f_j] \cap B$. Any parameter set sampled outside the hard boundaries is rejected to strictly enforce physical constraints. For convenience, the sampling range f_j is set to the same value for different reactions in the paper.

To ensure sample quality and reduce computational cost from numerical simulations, kinetic parameter combinations are filtered based on DNN-identified deviations that exceed our exploration range. Drawing inspiration from the principles of rejection sampling methodology,

the error tolerance at iteration t is governed by a threshold hyperparameter Θ_Q . A sample is retained only if its loss function satisfies $\tilde{\mathcal{L}}_Q(A) < \Theta_Q$ for all quantities of interest $Q \in \text{QoIs}$. Following filtration, the QoIs of retained samples are computed via numerical simulation using Cantera [35] and incorporated into an updated dataset for neural network training. As a heuristic rule, the sampling threshold monotonically decreases with iteration t . We select the current best model A^* as the temporary result of iteration t . The process terminates when the error function $\mathcal{L}(A^*)$ shows negligible improvement. In this work, we set the termination condition as when the difference in relative error between two consecutive iterations is less than 2%, ensuring stable convergence without excessive computation. The relative error will be calculated by $|\frac{Q(A^*) - Q^b}{Q^b}|$. We will also report the logarithmic error $|\log Q(A^*) - \log Q^b|$ for IDT. To better illustrate our method, we provide its pseudo-code in Alg. 1.

Fig. 1(d) visualizes the optimization process using two complementary subfigures in a coordinate system defined by two representative pre-exponential factors from the kinetic model. The upper subfigure tracks the progressive enhancement of DNN accuracy across iterations, culminating in identifying the global minimum for model errors. The optimization trajectory appears as a continuous path through the parameter space. The lower subfigure depicts the spatial distribution of sampled data points along this optimization trajectory, revealing how the sampling strategy evolves as the algorithm converges toward optimal parameters. This visualization illustrates DeePMO's ability to concentrate computational resources on the most influential parameter subspaces efficiently.

Algorithm 1 DeePMO

Input: Original mechanism A_0 , the size of model pool N , the benchmarks Q^b , the sample range f_j and max sample range $f_{bound,j}$, the weight of QoIs ω_Q , threshold parameter Θ_Q and the DNN hyperparameters.

```

1:  $A^* \leftarrow A_0$ ,  $A^* = (a_1, \dots, a_{nr})$ ;
2: for  $i = 1 : n$  do
3:   // Step 1: Generate model pool
4:   while number of  $A < N$  do
5:     Generate  $N$  samples  $A$  uniformly in  $[\log a_j - f_j, \log a_j + f_j]$ ;
6:     Filter the samples which are out of max sample range  $B = [\log a_j^0 - f_{bound,j}, \log a_j^0 + f_{bound,j}]$ ;
7:     if Iteration  $i \geq 1$  then
8:       Predict the QoIs of  $A$  By previous neural network DNN;
9:       // DNN-assistant screening
10:      For all  $A \in A$ , if  $\tilde{\mathcal{L}}_Q(A) < \Theta_Q$ , then the sample is retained in model pool  $A$ .
11:    end if
12:  end while
13:  Use Cantera to generate the QoIs  $\mathcal{Q}(A)$  of model pool  $A$ .
14: end for
15: // Step 2: Train the DNN for surrogate model
16: Train the DNN with input  $A$  and the label  $\mathcal{Q}(A)$ .
17: // Step 3: Find out the best sample
18:  $A^* \leftarrow \arg \min_i \{\mathcal{L}(A_i) : i = 1, \dots, N\}$ , with the error function  $\mathcal{L}$  defined in Eq. (4).
```

Output: The final optimized model A^*

3. Result and discussion

This section demonstrates the effectiveness of DeePMO through comprehensive case studies. To clarify the results presented in our plots, we use the following consistent terminology:

- **Benchmark:** Refers to the ground-truth target values for each QoI, sourced from cited experiments or detailed Cantera simulations.

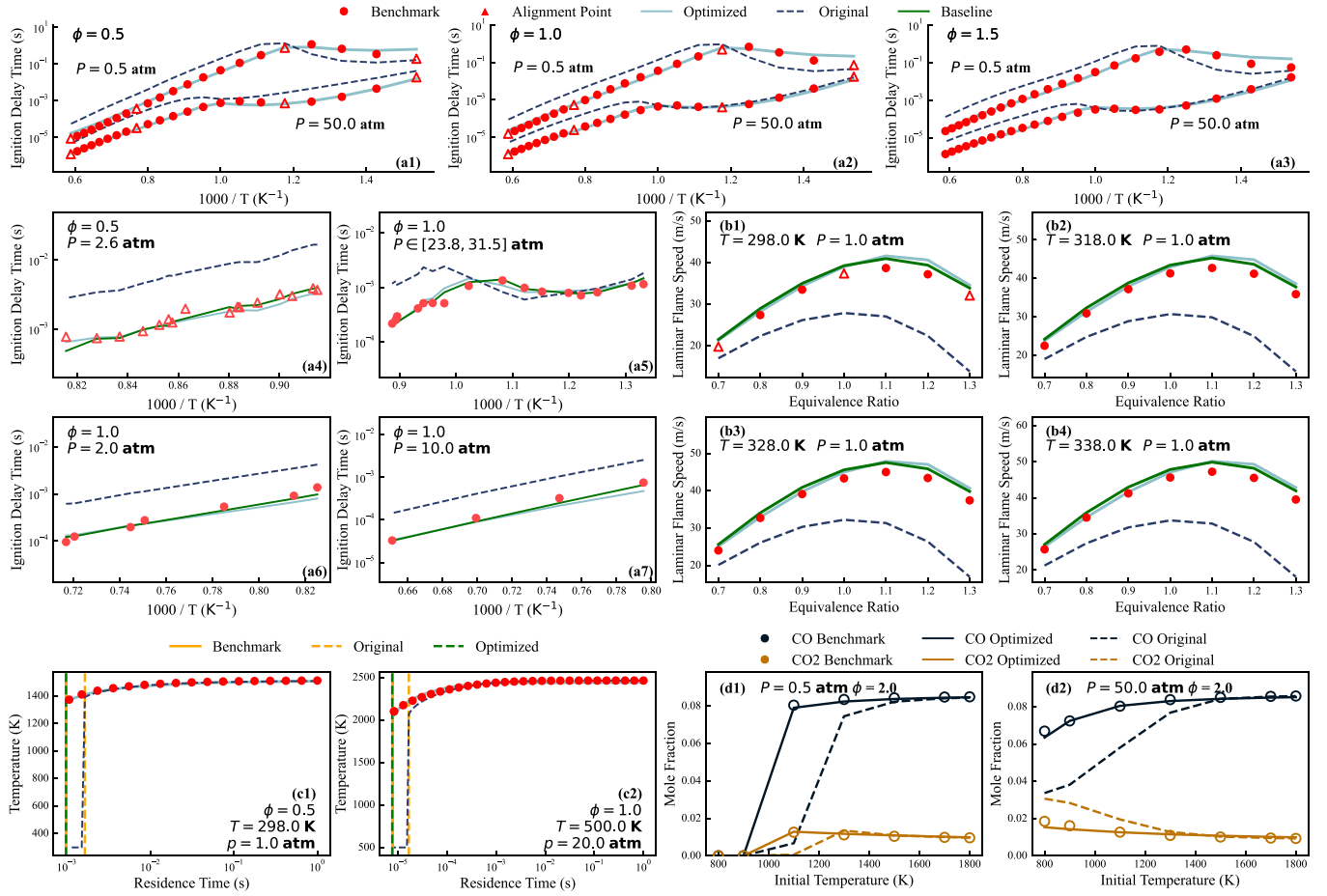


Fig. 2. Comparison between original and optimized n-heptane models: (a) IDT (b) LFS (c) PSRT(PSRex) and (d) Mole fraction in PSR predictions. The subplots (a1)–(a3) show the simulation results before and after the optimization, with average relative errors reduced from 285% (logarithmic difference: 0.53) to 19.48% (0.07). Subplots (a4)–(a7) utilizes the experiment data from [36–39] separately, where [36] is used for optimization and [37–39] for validation. We use the *baseline* from the detailed LLNL mechanism for reference. Subplots (b1)–(b4) use LFS measurements from [40]. The PSRT results come from simulation benchmarks where the PSR extinction time is highlighted. Subplot (d) shows the mole fraction of CO and CO₂ in PSR with residence time 0.05 s. We use the simulation data as benchmark.

- **Original:** Predictions from the initial, unoptimized kinetic model under benchmark conditions.
- **Alignment Points:** Subset of benchmark data used to compute the loss function $\mathcal{L}(A)$ in Eq. (4) during optimization.
- **Optimized:** Predictions from the final optimized model when evaluated on a set of validation conditions. These conditions were held out and not used during optimization. The validation set includes unseen conditions in the optimization process, allowing us to evaluate the model's robustness.

3.1. n-heptane model optimization using simulation and experimental data

We first examine a 27-species n-heptane model. It is derived from the detailed LLNL model (648 species, 4846 reactions) [41] and over-reduced by DeePMR [42,43] until 29 species, 145 reactions, and 152 pre-exponential factors. In this work, we optimize all of the pre-exponential factors. Although the reduced LLNL-29sp model (29 species, 145 reactions, 152 pre-exponential factors) preserves core predictive capabilities, over-simplification leads to degraded performance, posing a high-dimensional optimization challenge. We demonstrate DeePMO's efficacy in enhancing accuracy through parameter tuning, validated against benchmarks from detailed-model simulations and cited experiments [36], with extrapolation tests on independent datasets [37] confirming generalization and reproducibility.

The optimization targets multiple quantities of interest (QoIs). The alignment points for each QoI are defined as follows:

1. **IDT:** A total of 61 points are used, comprising 16 experimental conditions from Campbell et al. [36] and 45 simulation benchmarks from the detailed mechanism. The simulation conditions are uniformly distributed across temperatures $T \in \{750, 850, 950, 1050, 1700\}$ K, pressures $P \in \{0.5, 1, 50\}$ atm, and equivalence ratios $\phi \in \{0.5, 1, 2\}$.
2. **LFS:** 6 experimental points from Sileghem et al. [40] are used, covering $T = 298$ K, $P \in \{1, 10\}$ atm, and $\phi \in \{0.6, 1, 1.4\}$.
3. **PSRex/PSRT:** 18 points are defined by the combination of $T \in \{298, 470\}$ K, $P \in \{1, 20, 50\}$ atm, and $\phi \in \{0.5, 1, 2\}$.
4. **Mole Fractions in PSR:** 45 simulation benchmarks are used for the profiles of CO and CO₂. The conditions cover initial temperatures of $T \in \{900, 1100, 1300, 1500, 1700\}$ K, pressures of $P \in \{0.5, 1, 50\}$ atm, and equivalence ratios of $\phi \in \{0.5, 1, 2\}$.

For every iteration, we sample $N = 80,000$ to construct the model pool, and the sample range is set as $f_j = 0.1$. The hyperparameter weight w_Q is set as 1, 1, 0.8, 50 for IDT, PSR, LFS and mole fraction in PSR, respectively, to account for the disparate scales and relative importance of these QoIs, as informed by multi-objective balancing techniques in [24,25]. The DNN is trained with optimizer Adam [44] with learning rate $\text{lr} = 2\text{e-}5$ and L^2 regularization coefficient $\lambda =$

1e–2. The batch size *bs* is set to 2000, and the architecture scale is 3000, 2000, 2000. We strategically adopted an over-parameterized network design, as evidenced by the training phase consuming less computational budget than sampling operations. The hyperparameters selection and costs are discussed in Section 3.5.

Fig. 2 compares the performance of original and optimized models: (a) IDT predictions, (b) LFS predictions, (c) PSRT distributions, and (d) mole fractions of CO and CO₂ in PSR versus initial temperature. As the figures (a1) to (a3) illustrate, the optimization significantly improves model accuracy, reducing the average relative error in IDTs from 285% (0.53) to 19.48% (0.07) on the alignment simulation benchmarks. The DeePMO-optimized model generalizes well across a broad range of operating conditions, even when optimized on a relatively small set of alignment points. The model's accuracy is particularly high at high temperatures. Within the NTC region, a higher density of alignment points was employed to capture the sharp variations in IDT. Nevertheless, minor deviations in the low temperature regime between some test points and the benchmark data are still observed. Reflecting a trade-off between computational cost and accuracy, no more alignment points were added. Overall, DeePMO exhibits excellent **interpolation** capabilities for the IDT predictions.

To evaluate the extrapolation capability of the present model, it was further tested against experimental data from Zhang et al. [37], Shao et al. [38], and Liang et al. [39]. For comparison, a baseline curve, representing the detailed mechanism, is included to demonstrate its consistency with this set of experimental data. The results indicate that our optimized mechanism exhibits consistently excellent performance on the **extrapolation** datasets, including enhanced agreement with experimental LFS measurements. As specifically shown in Fig. 2(a5), in the negative temperature coefficient (NTC) region, our optimized model successfully corrects the inaccuracies of the reduced mechanism and accurately reproduces the NTC phenomena observed in the experiments.

As for the optimization of Laminar Flame Speed (LFS), we follow the LFS measurements from [40]. Given that the computational cost of LFS is significantly higher than that of other QoIs, selecting a small yet representative set of alignment points is preferable. An optimal approach is to choose three points at a fixed temperature and pressure, corresponding to lean ($\phi = 0.5$), stoichiometric ($\phi = 1.0$), and rich ($\phi = 1.5$) equivalence ratios. While the reduced mechanism initially underestimates the LFS values, all test points demonstrate satisfactory optimization results after DeePMO.

The final temperature of a Perfectly Stirred Reactor (PSR) is strongly correlated with its extinction limit. As indicated in Fig. 2(c), with residence time decreasing, the final temperature shows a slight decline before dropping abruptly at the extinction point. The reduced mechanism predicts a longer residence time at extinction compared to the detailed mechanism, which serves as the benchmark. This discrepancy is likely due to the slower overall reaction rate in the reduced model, caused by the omission of numerous reaction pathways. Our method successfully aligns the extinction limit with the benchmark while maintaining high accuracy for the PSRT.

For the mole fraction in PSR as the QoI, we follow the settings from [12], where a mixture of 10% fuel/oxidizer blend and 90% AR is fed into a PSR at various inlet temperatures. The residence time is fixed at 0.05 s to measure the resulting mole fractions of CO and CO₂. The reduced mechanism shows good agreement with the detailed mechanism under lean conditions ($\phi \leq 1$), but exhibits significant prediction discrepancies under rich conditions ($\phi > 1$), especially at lower temperatures. The proposed method successfully preserves the performance under both lean and stoichiometric conditions and enhances the prediction accuracy of species concentrations under rich conditions. Additionally, in recognition of the exponential nature of pressure dependencies, we conducted a further test to validate DeePMO's performance at pressure extremes. By extending the pressure range to include 0.1 atm and 100 atm, we found that the optimization capability of the method remains robust with the reduction of logarithmic differences from 0.40 to 0.09. More results and analysis of LLNL-29sp are detailed in Appendix.

3.2. Parameter optimization for ammonia and ammonia/alcohol models

We apply DeePMO to complex fuel systems: a detailed ammonia model [47] and an ammonia/methanol/ethanol model [48]. These applications demonstrate the method's ability to handle multi-component fuel systems with varying numbers of species. Through ablation experiments, we analyze the critical role of DNN in data sampling and parameter phase space exploration. Third, we validate DeePMO's effectiveness through extensive case studies spanning from methane to iso-octane models.

For pure ammonia combustion, we employ a mechanism developed by Zhan et al. [47], comprising 40 species, 272 reactions, and 289 adjustable parameters. IDT and LFS were selected as QoIs, with experimental data serving as benchmarks for optimization. The experimental data included IDT measurements from Chen et al. [45] and LFS measurements from Mei et al. [46]. For LFS, we used experimental data with an oxidizer composition of O₂/N₂ = 35/65, while model performance was additionally evaluated under air oxidizer conditions. We utilize 11 IDT alignment data and 6 LFS alignment data in total.

Fig. 3 presents the optimization results of the Zhan-2024 model. Compared to the original model, the optimized model significantly improves performance across all tested conditions. The average relative error decreases from 56.76% (0.36) to 14.19% (0.06) for IDT and from 12.84% to 4.23% for LFS on alignment points. We further tested the performance of this method on the dataset from Mei et al. [46] and observed a significant improvement. Additionally, we attempted to test the model's LFS performance for ammonia-air mixtures. Although the results were not entirely satisfactory, they still represented an improvement over the original mechanism. This indicates limitations in the method's extrapolation capability.

These improvements demonstrate DeePMO's effectiveness in optimizing ammonia combustion models. Unlike the reduced n-heptane (LLNL) model case in Section 3.1, this detailed ammonia model is larger and begins with smaller initial errors. Nevertheless, the significant improvements in both IDT and LFS predictions validate DeePMO's effectiveness.

To further evaluate DeePMO's capability on more complex mechanisms, we tested it on the more complex CEU–NH₃ model developed by Wang et al. [48] for ammonia/CH₃OH/CH₃CH₂OH/air mixture, comprising 91 species, 444 reactions and 478 adjustable parameters, representing a complex C–H–O–N system. IDT was selected as the QoI for optimization, given the simulation efficiency. Experimental data from an extensive set of studies [49–51] were used, from which 44 points were selected for training.

Fig. 4(a–c) demonstrates DeePMO's optimization results for IDT. The optimized CEU–NH₃ model shows significant improvements, with the average relative error reducing from 144.19% (0.38) to 28.56% (0.13) on alignment points. For NH₃–CH₃OH blending cases shown in (a), DeePMO substantially decreased average relative error from 166.80% (0.42) to 32.25% (0.15). The inclusion of NH₃–H₂ experimental data, which was important but omitted in the original model, led to dramatic improvements with just five experimental data points, reducing IDT average relative error from 1796.28% (1.2) to 39.37% (0.18) shown in (b). Notably, the optimization maintained the original model's accurate performance for pure ammonia IDT predictions without degradation, as shown in (c).

However, the results in Fig. 4(b1–b2) also highlight a limitation. As noted, while the model perfectly captures the single alignment point in each set, it fails to capture the overall trend and endpoints. We concur that this is a direct consequence of using sparse alignment data; a single point provides insufficient information to constrain the slope of the IDT curve during optimization.

In context, the original CEU–NH₃ model was not designed for NH₃–H₂ conditions, explaining its poor baseline. DeePMO's error reduction remains a notable advance, yet this case emphasizes the need to bolster extrapolability. Improvements could include strategic alignment

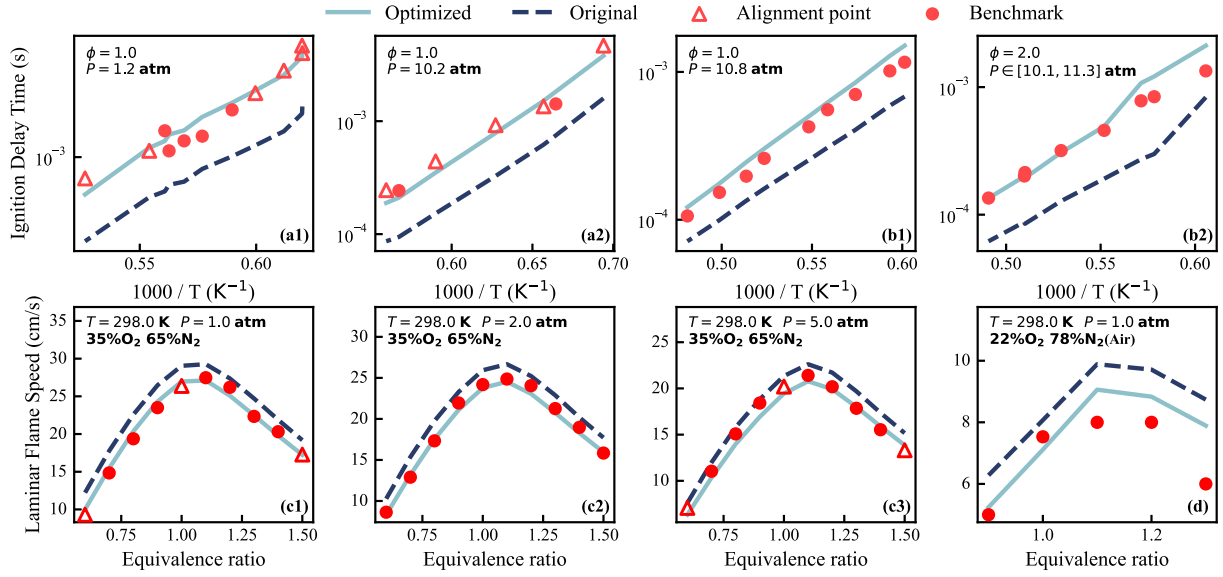


Fig. 3. Comparison between original and optimized ammonia models (Zhan-2024). (a) Ignition delay time versus initial temperature, the experiment data of Chen et al. [45] is selected as the benchmark and partially for training; (b) Ignition delay time validation on the experiment data from Mei et al. [46], which is excluded from optimization. (c) Laminar flame speed versus equivalence ratio in oxygen-rich conditions, and (d) Laminar flame speed in air. All LFS experiment measurements are from Mei et al. [46].

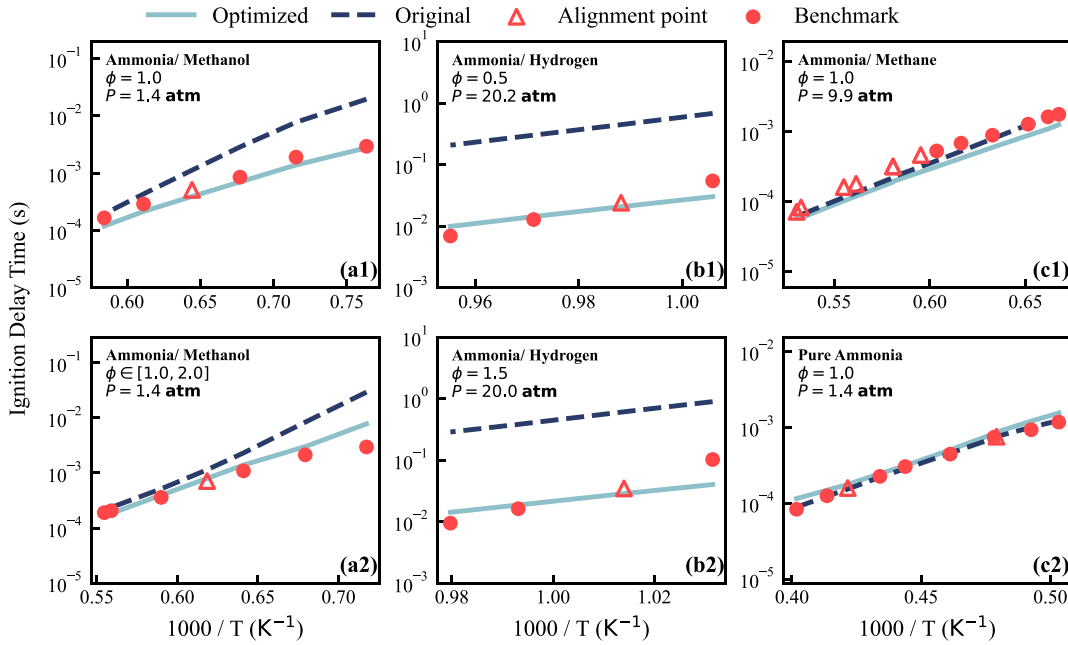


Fig. 4. Comparison of IDT before and after optimization for CEU-NH₃. Figs. (a–c) show the variation of IDT with temperature, pressure, and blending ratio under different types of combustible substances. The primary difference among these three figures is the type of fuel burned. (a): NH₃+CH₃OH. (b): NH₃+H₂. (c): NH₃.

point selection – e.g., two or three points spanning the temperature range (start, middle, end) – to better inform the optimizer of system dynamics. Future work may also integrate physics-informed priors or multi-fidelity data for enhanced robustness across unseen conditions.

The iterative sampling-training-inference process is fundamental to the DeePMO scheme, with the DNN serving as the cornerstone for rapid screening and global minimum identification. Fig. 5 presents an ablation study comparing performance with and without DNN assistance during iterative evolution, and we select the mechanism from Zhan et al. [47] as the target. Without DNN screening, both median and minimum relative errors remain nearly constant across iterations, indicating negligible optimization capability. In contrast, with DNN-guided sample selection, the error decreases substantially from nearly

60% to below 30% within four iterations, ultimately converging to 20% (IDT 14. 19% + LFS 4. 23%). This improvement demonstrates the effectiveness of DNN-based sample selection in enhancing training sample quality. The DNN screening enables efficient exploration of high-dimensional parameter space, facilitating optimal parameter combination identification.

3.3. Extensive validation across diverse fuel models

This section demonstrates the broad applicability of DeePMO in various combustion models. We validated the method using diverse fuel types: alkane models from LLNL, including C₂H₄, C₄H₁₀, iso-octane [52], n-heptane model from [52], alcohols such as n-pentanol

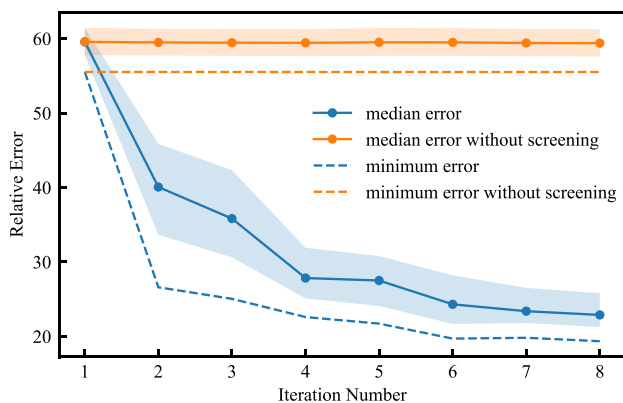


Fig. 5. Ablation analysis using ammonia model [47] regarding IDT and LFS as QoIs. The solid line represents the median performance of dataset $A^{(i)}$ at each iteration i , while the dashed line shows the minimum error sample performance A^* . The shaded areas indicate the quartile ranges of performance distribution.

[53], and different models for carbon-free fuel ammonia [54–56]. The alkane/alcohols models are reduced by DeePMR and then optimized by DeePMO on IDT with the detailed mechanism LLNL as the benchmark. The ammonia models are detailed mechanisms, and we apply IDT experiment data as the benchmark from Chen et al. [45]. As demonstrated in Fig. 6, DeePMO significantly improves the accuracy of these models across different scenarios. Empirically, these experiments show robustness under different hyperparameter settings, and we provide a detailed analysis in Section 3.4.

These results demonstrate that DeePMO exhibits robust optimization capabilities across a broad spectrum of combustion mechanisms, whether over reduced models or detailed mechanisms against simulation or experiment datasets. For complex reaction frameworks, the optimized relative errors can be constrained within 20%, aligning with the experimental measurement uncertainty of IDT under standard conditions. The relative errors can be further reduced to lower levels for mechanisms with simplified architectures.

3.4. Ablation experiment of hyperparameters

This section details the ablation studies conducted on the hyperparameters of the DeePMO methodology, focusing on critical parameters such as network architecture, batch size, learning rate, and the parameter sampling range. These systematic investigations underscore the robustness and effectiveness of the DeePMO algorithm across diverse settings. Guided by insights from prior studies on ANN-based parameter optimization [24,25,31], we established a baseline configuration for our analysis with the following specifications: a network architecture of [3000, 2000, 2000], a batch size of 1000, a learning rate of $1e-5$, and a sampling range of 0.15.

We employed a controlled variable methodology for our ablation studies, systematically varying individual parameters while keeping others constant. We evaluated multiple experimental configurations within prescribed hyperparameter adjustment ranges and reported peak performance metrics from the optimization outcomes. Detailed experimental findings are presented in Table 1. The ablation study reveals two critical insights regarding the DeePMO algorithm's characteristics: (1) DeePMO demonstrates stable optimization performance across a rational hyperparameter range, exhibiting low sensitivity to architecture-related hyperparameters while showing pronounced responsiveness to the sampling range f_j . (2) The original network architecture contains substantial redundancy—preliminary experiments demonstrate that a 90% scaled-down network (100-50-50 vs. 3000-2000-2000) maintains

Table 1

Hyperparameter ablation studies on DeePMO of mechanism [47]. All results represent the best relative error of IDT and LFS under the corresponding hyperparameter configurations. For the baseline setting, we have $f_j = 0.15$, $lr = 1e-5$, $bs = 1000$, $scale = 3000, 2000, 2000$.

Ablation studies		IDT (%)	LFS (%)
Baseline		13.92	5.87
Sampling range f_j	0.1	14.62	5.74
	0.2	12.95	5.48
	0.25	14.64	5.02
	0.3	12.66	2.34
	0.35	11.80	3.22
	0.4	12.67	6.51
	0.45	13.23	1.82
	0.5	Method diverged	
Learning rate lr	2e–6	14.95	5.26
	5e–5	14.19	6.16
Batch size bs	500	13.66	5.79
	2000	14.38	6.56
Architecture scale scale	100, 50, 50	14.27	5.85
	50, 10, 10	Method diverged	

comparable accuracy. This indicates that both the scale and computational cost can be significantly reduced without compromising performance, suggesting that optimization outcomes are not sensitive to neural network sizes provided the width remains reasonably high.

Similarly with the mechanism [47], we conducted an ablation study on the CEU-NH₃ mechanism (shown in Table 2), focusing specifically on the ‘sampling range’ hyperparameter f_j . The results, compiled after nine optimization iterations, indicate that narrower sampling ranges tend to be less efficient. Conversely, while larger sampling ranges offer potential for accelerated optimization progress, they may introduce divergence challenges that compromise stability. More discussion and about role of sampling range f_j could be found in Appendix.

3.5. Computational cost breakdown

The overall computational cost comprises two principal components: (1) Screening and simulation, (2) DNN optimization. The former constitutes mechanism-dependent costs primarily governed by chemical complexity and sampling density requirements, while the latter reflects hardware-accelerated training overhead. With the n-heptane LLNL-29sp experiment as an example, our result reveals:

Per-sample simulation costs:

1. Ignition delay time (IDT): 80 s/sample (61 thermodynamic conditions)
2. Extinction time/temperature profile in Perfectly stirred reactor (PSRex/PSRT): 16 s/sample (27 flow configurations)
3. Laminar flame speed (LFS): 52 s/sample (6 alignment points in 3 pressure regimes)
4. Mole fraction in PSR: 8 s/sample (45 alignment points in 3 pressure regimes)

Large-scale parallelization: 80,000-sample campaigns require ≈ 2 h 30 min using 1280 CPU cores by parallel calculation.

Neural network training: 700-epoch convergence was achieved in 45 min using a single NVIDIA GeForce RTX 2080 Ti with network dimensions (scale) set to 3000, 2000, 2000. This time cost decreased to approximately 20 min when reducing network dimensions to 100, 50, 50.

Neural network training represents a relatively minor component of the overall computational demand compared to sampling processes, which utilize thousands of CPU cores. The simulation time for LFS calculations proved to be the most computationally intensive component, making it prohibitively expensive for direct large-scale mechanism optimization (see Table 2).

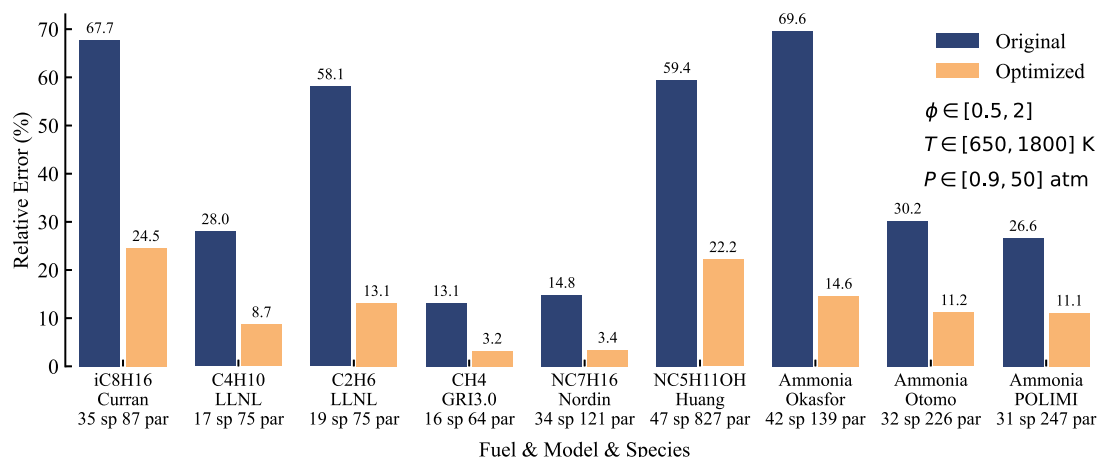


Fig. 6. Comparison of IDT relative errors before and after DeePMO optimization across various combustion models on alignment points. For the alkane mechanisms, a reduction was performed using the DeePMR method prior to optimization. The n-pentanol and other ammonia mechanisms were optimized directly from the detailed mechanisms without simplification. The number of species and the total optimized parameters are labeled in the figure.

Table 2

Hyperparameter Ablation Studies on DeePMO of CEU-NH₃ [48]. All results represent the best relative error of IDT and LFS under the sampling range. For all settings, we have $\text{lr} = 1\text{e-}5$, $\text{bs} = 1000$, $\text{scale} = 3000, 2000, 2000$.

Ablation studies	IDT (%)
0.05 (baseline)	28.56
0.2	28.94
0.25	23.29
0.3	25.98
0.35	17.19
0.4	20.46
0.5	Method diverged

4. Conclusion

This work presents DeePMO, a novel deep learning-based approach for optimizing chemical kinetic models. Through extensive validation across diverse fuel models, we have demonstrated DeePMO's effectiveness in improving model accuracy while maintaining computational efficiency. The method successfully reduced average relative errors in various quantities of interest across all tested models.

DeePMO's key innovations – its iterative sampling-learning-inference strategy and hybrid DNN architecture – effectively tackle core challenges in data-driven chemical kinetics optimization, including high-dimensional spaces and heterogeneous metrics. This unified framework enables robust handling of diverse QoIs from simulations and experiments, with rigorous enforcement of physical uncertainty bounds ($f_{\text{bound},j}$) and tunable sampling ranges (f_j) for plausible, efficient exploration. Ablation studies confirm DNN-guided sampling's pivotal role in enhancing convergence and sample efficiency, while extrapolation validations on independent datasets affirm generalization.

DeePMO's successful application to various fuel models demonstrates its versatility and robustness, including methane, ethane, butane, n-heptane, n-pentanol, ammonia, and their mixtures. Particularly noteworthy is its performance in optimizing complex models like CEU-NH₃, where it reduced IDT relative errors from >100% to 20% (logarithmic differences from 0.45 to 0.09) using sparse alignment points, with extrapolation validations on independent datasets confirming generalization without overfitting. The method's capability to incorporate both experimental measurements and simulation data further enhances its practical utility.

These results establish DeePMO as a powerful, scalable tool for combustion model optimization, with demonstrated robustness across high-dimensional mechanisms and diverse QoIs. Limitations include

potential challenges in capturing trends with highly sparse alignment data, as observed in single-point cases, underscoring the need for strategic QoI selection to ensure well-posed problems and uncorrelated constraints. Future work could extend the method to additional combustion characteristics, more complex fuel mixtures, and integrated algorithms for optimal data point curation, further broadening its applicability in combustion chemistry research.

CRediT authorship contribution statement

Pengxiao Lin: Writing – review & editing, Writing – original draft, Visualization, Methodology, Investigation, Formal analysis, Data curation, Conceptualization. **Yuntian Zhou:** Writing – original draft, Investigation, Data curation. **Zhiwei Wang:** Data curation, Conceptualization. **Weizong Wang:** Supervision, Resources, Funding acquisition. **Zheng Chen:** Supervision, Resources, Funding acquisition. **Zhi-Qin John Xu:** Writing – review & editing, Writing – original draft, Supervision, Resources, Methodology, Formal analysis. **Tianhan Zhang:** Writing – review & editing, Supervision, Formal analysis, Data curation, Conceptualization.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Acknowledgments

This work is sponsored by the National Natural Science Foundation of China Grant No. 92470127, 12371511, 92270203, the National Key R&D Program of China Grant No. 2019YFA0709503, and the Open Project of the National Key Laboratory of Scramjet Technology No. WDZC6142703202403.

Appendix A. Supplementary data

Supplementary material related to this article can be found online at <https://doi.org/10.1016/j.jaecs.2025.100402>.

Data availability

Data will be made available on request.

References

- [1] Frenklach M, Wang H, Rabinowitz MJ. Optimization and analysis of large chemical kinetic mechanisms using the solution mapping method—combustion of methane. *Prog Energy Combust Sci* 1992;18(1):47–73. [http://dx.doi.org/10.1016/0360-1285\(92\)90032-V](http://dx.doi.org/10.1016/0360-1285(92)90032-V), URL <https://www.sciencedirect.com/science/article/pii/036012859290032V>.
- [2] Zhang T, Susa AJ, Hanson RK, Ju Y. Studies of the dynamics of autoignition assisted outwardly propagating spherical cool and double flames under shock-tube conditions. *Proc Combust Inst* 2020;1–9. <http://dx.doi.org/10.1016/j.proci.2020.06.089>, URL <https://linkinghub.elsevier.com/retrieve/pii/S1540748920301449>. GSCC: 0000028 Publisher: Elsevier Inc.
- [3] Zhang T, Susa AJ, Hanson RK, Ju Y. Two-dimensional simulation of cool and double flame formation induced by the laser ignition under shock-tube conditions. *Proc Combust Inst* 2022;39:4. <http://dx.doi.org/10.1016/j.proci.2022.08.068>, URL <https://linkinghub.elsevier.com/retrieve/pii/S1540748922003492>. GSCC: 0000012.
- [4] Wang H, Sheen DA. Combustion kinetic model uncertainty quantification, propagation and minimization. *Prog Energy Combust Sci* 2015;47:1–31. <http://dx.doi.org/10.1016/j.peccs.2014.10.002>.
- [5] Elliott L, Ingham D, Kyne A, Mera N, Pourkashanian M, Wilson C. Genetic algorithms for optimisation of chemical kinetics reaction mechanisms. *Prog Energy Combust Sci* 2004;30(3):297–328. <http://dx.doi.org/10.1016/j.peccs.2004.02.002>, URL <https://www.sciencedirect.com/science/article/pii/S0360128504000115>.
- [6] Kelly M, Fortune M, Bourque G, Dooley S. Machine learned compact kinetic models for methane combustion. *Combust Flame* 2023;253:112755. <http://dx.doi.org/10.1016/j.combustflame.2023.112755>.
- [7] Polifke W, Geng W, Döbbling K. Optimization of rate coefficients for simplified reaction mechanisms with genetic algorithms. *Combust Flame* 1998;113(1):119–34. [http://dx.doi.org/10.1016/S0010-2180\(97\)00212-5](http://dx.doi.org/10.1016/S0010-2180(97)00212-5), URL <https://www.sciencedirect.com/science/article/pii/S0010218097002125>.
- [8] Harris S, Elliott L, Ingham D, Pourkashanian M, Wilson C. The optimisation of reaction rate parameters for chemical kinetic modelling of combustion using genetic algorithms. *Comput Methods Appl Mech Engrg* 2000;190(8):1065–90. [http://dx.doi.org/10.1016/S0045-7825\(99\)00466-1](http://dx.doi.org/10.1016/S0045-7825(99)00466-1), URL <https://www.sciencedirect.com/science/article/pii/S0045782599004661>.
- [9] Kim K, Wiersema PW, Ryu JI, Mayhew E, Temme J, Kweon C-B, Lee T. Data-driven approaches to optimize chemical kinetic models. In: AIAA SCITECH 2022 forum. AIAA sciTech forum, American Institute of Aeronautics and Astronautics; 2021. <http://dx.doi.org/10.2514/6.2022-0225>.
- [10] Li G, Rosenthal C, Rabitz H. High dimensional model representations. *J Phys Chem* 2001;105(33):7765–77. <http://dx.doi.org/10.1021/jp010450t>.
- [11] Ziehn T, Tomlin AS. A global sensitivity study of sulfur chemistry in a premixed methane flame model using HDMR. *Int J Chem Kinet* 2008;40(11):742–53. <http://dx.doi.org/10.1002/kin.20367>.
- [12] Fürst M, Bertolino A, Cuoci A, Faravelli T, Frassoldati A, Parente A. OptiSMOKE++: A toolbox for optimization of chemical kinetic mechanisms. *Comput Phys Comm* 2021;264:107940. <http://dx.doi.org/10.1016/j.cpc.2021.107940>, URL <https://www.sciencedirect.com/science/article/pii/S0010465521000680>.
- [13] Papp M, Varga T, Busai Á, Zsély IG, Nagy T, Turányi T. Optima++ v2.1: A general C++ framework for performing combustion simulations and mechanism optimization. 2021.
- [14] Turányi T, Nagy T, Zsély IG, Cserháti M, Varga T, Szabó BT, Sedyó I, Kiss PT, Zempléni A, Curran HJ. Determination of rate parameters based on direct and indirect measurements. *Int J Chem Kinet* 2012;44(5):284–302. <http://dx.doi.org/10.1002/kin.20717>, arXiv:<https://onlinelibrary.wiley.com/doi/pdf/10.1002/kin.20717>. URL <https://onlinelibrary.wiley.com/doi/abs/10.1002/kin.20717>.
- [15] Goitom SK, Papp M, Kovács M, Nagy T, Zsély IG, Turányi T, Pál L. Efficient numerical methods for the optimisation of large kinetic reaction mechanisms. *Combust Theory Model* 2022;26(6):1071–97. <http://dx.doi.org/10.1080/13647830.2022.2110945>, arXiv:<https://doi.org/10.1080/13647830.2022.2110945>.
- [16] Seo J, Kim S, Jalalvand A, Conlin R, Rothstein A, Abbate J, Erickson K, Wai J, Shousha R, Kolemen E. Avoiding fusion plasma tearing instability with deep reinforcement learning. *Nature* 2024. <http://dx.doi.org/10.1038/s41586-024-07024-9>.
- [17] Jumper J, Evans R, Pritzl A, Green T, Figurnov M, Ronneberger O, Tunyasuvunakool K, Bates R, Židek A, Potapenko A, Bridgland A, Meyer C, Kohl SAA, Ballard AJ, Cowie A, Romero-Paredes B, Nikolov S, Jain R, Adler J, Back T, Petersen S, Reiman D, Clancy E, Zielinski M, Steinegger M, Pacholska M, Berghammer T, Bodenstern S, Silver D, Vinyals O, Senior AW, Kavukcuoglu K, Kohli P, Hassabis D. Highly accurate protein structure prediction with AlphaFold. *Nature* 2021;596(7873):583–9. <http://dx.doi.org/10.1038/s41586-021-03819-2>.
- [18] Zhang T, Zhang Y, E W, Ju Y. DLODE: a deep learning-based ODE solver for chemistry kinetics. In: AIAA scitech 2021 forum. American Institute of Aeronautics and Astronautics; 2021. <http://dx.doi.org/10.2514/6.2021-1139>, GSCC: 0000011. arXiv:2012.12654. URL <http://arxiv.org/abs/2012.12654>.
- [19] Zhang T, Yi Y, Xu Y, Chen ZX, Zhang Y, E W, Xu Z-QJ. A multi-scale sampling method for accurate and robust deep neural network to predict combustion chemical kinetics. *Combust Flame* 2022;245:112319. <http://dx.doi.org/10.1016/j.combustflame.2022.112319>, URL <https://linkinghub.elsevier.com/retrieve/pii/S0010218022003340>. GSCC: 0000069 TLD: The results reveal that the DNN trained by the manifold data can capture the chemical kinetics in limited configurations but cannot remain robust toward perturbation, which is inevitable for the DNN coupled with the flow field.
- [20] Wang T, Yi Y, Yao J, Xu Z-QJ, Zhang T, Chen Z. Enforcing physical conservation in neural network surrogate models for complex chemical kinetics. *Combust Flame* 2025;275:114105. <http://dx.doi.org/10.1016/j.combustflame.2025.114105>, URL <https://linkinghub.elsevier.com/retrieve/pii/S0010218025001439>. GSCC: 0000003 TLD: This work proposes a novel ANN approach with hard physical constraints (ANN-hard) for chemical source term calculations that strictly enforce conservation laws (mass, energy, and element) and demonstrates superior stability and physical accuracy by preventing error accumulation.
- [21] Yao J, Yi Y, Hang L, Wang W, Zhang Y, Zhang T, Xu Z-QJ, et al. Solving multiscale dynamical systems by deep learning. *Comput Phys Comm* 2025;109802.
- [22] Zhang X, Yi Y, Wang L, Xu Z-QJ, Zhang T, Zhou Y. Deep neural networks for modeling astrophysical nuclear reacting flows. 2025, arXiv preprint arXiv:2504.14180.
- [23] LeCun Y, Bengio Y, Hinton G. Deep learning. *Nature* 2015;521(7553):436–44. <http://dx.doi.org/10.1038/nature14539>.
- [24] Ji W, Su X, Pang B, Li Y, Ren Z, Deng S. SGD-based optimization in modeling combustion kinetics: Case studies in tuning mechanistic and hybrid kinetic models. *Fuel* 2022;324:124560. <http://dx.doi.org/10.1016/j.fuel.2022.124560>.
- [25] Oh J-H, Wiersema P, Kim K, Mayhew E, Temme J, Kweon C-B, Lee T. Fast uncertainty reduction of chemical kinetic models with complex spaces using hybrid response-surface networks. *Combust Flame* 2023;253:112772. <http://dx.doi.org/10.1016/j.combustflame.2023.112772>.
- [26] Li S, Yang B, Qi F. Accelerate global sensitivity analysis using artificial neural network algorithm: Case studies for combustion kinetic model. *Combust Flame* 2016;168:53–64. <http://dx.doi.org/10.1016/j.combustflame.2016.03.028>.
- [27] Su X, Ji W, An J, Ren Z, Deng S, Law CK. Kinetics parameter optimization of hydrocarbon fuels via neural ordinary differential equations. *Combust Flame* 2023;251:112732. <http://dx.doi.org/10.1016/j.combustflame.2023.112732>.
- [28] Owoyele O, Pal P. ChemNODE: A neural ordinary differential equations framework for efficient chemical kinetic solvers. *Energy AI* 2022;7:100118. <http://dx.doi.org/10.1016/j.egyai.2021.100118>, URL <https://www.sciencedirect.com/science/article/pii/S2666546821000677>.
- [29] Fedorov A, Perechodjuk A, Linke D. Kinetics-constrained neural ordinary differential equations: Artificial neural network models tailored for small data to boost kinetic model development. *Chem Eng J* 2023;477:146869. <http://dx.doi.org/10.1016/j.ccej.2023.146869>, URL <https://www.sciencedirect.com/science/article/pii/S1385894723056000>.
- [30] Zhang Y, Dong W, Vandewalle LA, Xu R, Smith GP, Wang H. Neural network approach to response surface development for reaction model optimization and uncertainty minimization. *Combust Flame* 2023;251:112679. <http://dx.doi.org/10.1016/j.combustflame.2023.112679>, URL <https://www.sciencedirect.com/science/article/pii/S0010218023000640>.
- [31] Wang Y, Liu C, Tao C, Law CK, Yang B. Efficient combustion kinetic parameter optimization via variational inference. *Proc Combust Inst* 2024;40(1):105550. <http://dx.doi.org/10.1016/j.proci.2024.105550>, URL <https://www.sciencedirect.com/science/article/pii/S1540748924003584>.
- [32] Chen H, Li Q, Deng S. Fast QoI-Oriented Bayesian experimental design with unified neural response surfaces for kinetic uncertainty reduction. *Energy Fuels* 2024;38(16):15630–41. <http://dx.doi.org/10.1021/acs.energyfuels.4c02299>.
- [33] Zhou Z, Lin K, Wang Y, Wang J, Law CK, Yang B. OptEx: An integrated framework for experimental design and combustion kinetic model optimization. *Combust Flame* 2022;245:112298. <http://dx.doi.org/10.1016/j.combustflame.2022.112298>, URL <https://www.sciencedirect.com/science/article/pii/S0010218022003133>.
- [34] Li M, Hu H, Lu L, Zhang H. Development of compact mechanism for lithium-ion battery venting gas fires using cantera ordinary differential equation neural network algorithm. *Appl Energy Combust Sci* 2025;22:100326. <http://dx.doi.org/10.1016/j.jaecs.2025.100326>, URL <https://www.sciencedirect.com/science/article/pii/S2666352X25000081>.
- [35] Goodwin DG, Speth RL, Moffat HK, Weber BW. Cantera: An object-oriented software toolkit for chemical kinetics, thermodynamics, and transport processes. 2021. <http://dx.doi.org/10.5281/zenodo.4527812>.
- [36] Campbell MF, Wang S, Davidson DF, Hanson RK. Shock tube study of normal heptane first-stage ignition near 3.5 Atm. *Combust Flame* 2018;198:376–92. <http://dx.doi.org/10.1016/j.combustflame.2018.08.008>.
- [37] Zhang D, Wang Y, Zhang C, Li P, Li X. Experimental and numerical investigation of vitiation effects on the auto-ignition of n-heptane at high temperatures. *Energy* 2019;174:922–31. <http://dx.doi.org/10.1016/j.energy.2019.03.035>, URL <https://www.sciencedirect.com/science/article/pii/S0360544219304396>.
- [38] Shao J, Choudhary R, Peng Y, Davidson DF, Hanson RK. A shock tube study of n-heptane, iso-octane, n-dodecane and iso-octane/n-dodecane blends oxidation at elevated pressures and intermediate temperatures. *Fuel* 2019;243:541–53. <http://dx.doi.org/10.1016/j.fuel.2019.01.152>, URL <https://www.sciencedirect.com/science/article/pii/S001623611930153X>.

- [39] Liang J, Zhang Z, Li G, Wan Q, Xu L, Fan S. Experimental and kinetic studies of ignition processes of the methane–n-heptane mixtures. *Fuel* 2019;235:522–9. <http://dx.doi.org/10.1016/j.fuel.2018.08.041>, URL <https://www.sciencedirect.com/science/article/pii/S001623611831411X>.
- [40] Sileghem L, Alekseev VA, Vancoillie J, Van Geem KM, Nilsson EJK, Verhelst S, Konnov AA. Laminar burning velocity of gasoline and the gasoline surrogate components iso-octane, n-heptane and toluene. *Fuel* 2013;112:355–65. <http://dx.doi.org/10.1016/j.fuel.2013.05.049>.
- [41] Mehl M, Pitz WJ, Westbrook CK, Curran HJ. Kinetic modeling of gasoline surrogate components and mixtures under engine conditions. *Proc Combust Inst* 2011;33(1):193–200. <http://dx.doi.org/10.1016/j.proci.2010.05.027>.
- [42] Wang Z, Zhang Y, Lin P, Zhao E, E W, Zhang T, Xu Z-QJ. Deep mechanism reduction (DeePMR) method for fuel chemical kinetics. *Combust Flame* 2024;261:113286. <http://dx.doi.org/10.1016/j.combustflame.2023.113286>.
- [43] Song Y, Shen W, Bai S, Li S, Liang X, Shao J, Liu Z, Feng G, Zhao C, He X, Li Y, Liang J, Guan X, Zhang T, Wang Z, Xu Z-QJ, Chen D, Wang K. Modeling combustion chemistry of China aviation kerosene (RP-3) through the Hy-Chem approach. *Combust Flame* 2025;280:114339. <http://dx.doi.org/10.1016/j.combustflame.2025.114339>, URL <https://linkinghub.elsevier.com/retrieve/pii/S0010218025003761>.
- [44] Kingma DP, Ba J. Adam: A method for stochastic optimization. 2017, [arXiv:1412.6980](https://arxiv.org/abs/1412.6980). URL <https://arxiv.org/abs/1412.6980>.
- [45] Chen J, Jiang X, Qin X, Huang Z. Effect of hydrogen blending on the high temperature auto-ignition of ammonia at elevated pressure. *Fuel* 2021;287:119563. <http://dx.doi.org/10.1016/j.fuel.2020.119563>.
- [46] Mei B, Zhang X, Ma S, Cui M, Guo H, Cao Z, Li Y. Experimental and kinetic modeling investigation on the laminar flame propagation of ammonia under oxygen enrichment and elevated pressure conditions. *Combust Flame* 2019;210:236–46. <http://dx.doi.org/10.1016/j.combustflame.2019.08.033>, URL <https://www.sciencedirect.com/science/article/pii/S0010218019303979>.
- [47] Zhan H, Li S, Yin G, Hu E, Huang Z. Experimental and kinetic study of ammonia oxidation and NOx emissions at elevated pressures. *Combust Flame* 2024;263:113129. <http://dx.doi.org/10.1016/j.combustflame.2023.113129>.
- [48] Wang Z, Han X, He Y, Zhu R, Zhu Y, Zhou Z, Cen K. Experimental and kinetic study on the laminar burning velocities of NH₃ mixing with CH₃OH and C₂H₅OH in premixed flames. *Combust Flame* 2021;229:111392.
- [49] He X, Shu B, Nascimento D, Moshhammer K, Costa M, Fernandes R. Auto-ignition kinetics of ammonia and ammonia/hydrogen mixtures at intermediate temperatures and high pressures. *Combust Flame* 2019;206:189–200. <http://dx.doi.org/10.1016/j.combustflame.2019.04.050>.
- [50] Li X, Ma Z, Jin Y, Wang X, Xi Z, Hu S, Chu X. Effect of methanol blending on the high-temperature auto-ignition of ammonia: An experimental and modeling study. *Fuel* 2023;339:126911. <http://dx.doi.org/10.1016/j.fuel.2022.126911>.
- [51] Mathieu O, Petersen EL. Experimental and modeling study on the high-temperature oxidation of ammonia and related NO_x chemistry. *Combust Flame* 2015;162(3):554–70. <http://dx.doi.org/10.1016/j.combustflame.2014.08.022>.
- [52] Curran HJ, Gaffuri P, Pitz WJ, Westbrook CK. A comprehensive modeling study of iso-octane oxidation. *Combust Flame* 2002;129(3):253–80. [http://dx.doi.org/10.1016/S0010-2180\(01\)00373-X](http://dx.doi.org/10.1016/S0010-2180(01)00373-X).
- [53] Huang H, Lv D, Zhu J, Chen Y, Zhu Z, Pan M, Huang R, Jia C. Development and validation of a new reduced diesel/n-pentanol mechanism for diesel engine applications. *Energy Fuels* 2018;32(9):9934–48. <http://dx.doi.org/10.1021/acs.energyfuels.8b02083>.
- [54] Okafor EC, Naito Y, Colson S, Ichikawa A, Kudo T, Hayakawa A, Kobayashi H. Experimental and numerical study of the laminar burning velocity of CH₄–NH₃–air premixed flames. *Combust Flame* 2018;187:185–98. <http://dx.doi.org/10.1016/j.combustflame.2017.09.002>.
- [55] Otomo J, Koshi M, Mitsumori T, Iwasaki H, Yamada K. Chemical kinetic modeling of ammonia oxidation with improved reaction mechanism for ammonia/air and ammonia/hydrogen/air combustion. *Int J Hydrog Energy* 2018;43(5):3004–14. <http://dx.doi.org/10.1016/j.ijhydene.2017.12.066>.
- [56] Stagni A, Cavallotti C, Arunthanayothin S, Song Y, Herbinet O, Battin-Leclerc F, Faravelli T. An experimental, theoretical and kinetic-modeling study of the gas-phase oxidation of ammonia. *React Chem Eng* 2020;5:696–711. <http://dx.doi.org/10.1039/C9RE00429G>.