

MACE: A Machine learning Approach to Chemistry EmulationSILKE MAES ¹, FREDERIK DE CEUSTER ¹, MARIE VAN DE SANDE ^{2,3} AND LEEN DECIN ^{1,4}¹*Institute of Astronomy, KU Leuven, Celestijnenlaan 200D, B-3001 Leuven, Belgium*²*Leiden Observatory, Leiden University, PO Box 9513, 2300 RA Leiden, The Netherlands*³*School of Physics and Astronomy, University of Leeds, Leeds LS2 9JT, United Kingdom*⁴*School of Chemistry, University of Leeds, Leeds LS2 9JT, United Kingdom***ABSTRACT**

The chemistry of an astrophysical environment is closely coupled to its dynamics, the latter often found to be complex. Hence, to properly model these environments a 3D context is necessary. However, solving chemical kinetics within a 3D hydro simulation is computationally infeasible for even a modest parameter study. In order to develop a feasible 3D hydro-chemical simulation, the classical chemical approach needs to be replaced by a faster alternative. We present MACE, a Machine learning Approach to Chemistry Emulation, as a proof-of-concept work on emulating chemistry in a dynamical environment. Using the context of AGB outflows, we have developed an architecture that combines the use of an autoencoder (to reduce the dimensionality of the chemical network) and a set of latent ordinary differential equations (that are solved to perform the temporal evolution of the reduced features). Training this architecture with an integrated scheme makes it possible to successfully reproduce a full chemical pathway in a dynamical environment. MACE outperforms its classical analogue on average by a factor 26. Furthermore, its efficient implementation in PyTorch results in a sub-linear scaling with respect to the number of hydrodynamical simulation particles.

Keywords: Astrochemistry (75) – Computational methods (1965) – Astronomy software (1855) – Chemical reaction network models (2237) – Asymptotic giant branch stars (2100) – Stellar winds (1636)

1. INTRODUCTION

Astrochemistry, the study of chemistry in space, is a powerful tool. Combining observations of chemical species in astrophysical objects with theoretical predictions allows us to study the physical conditions, as well as to estimate its chemical composition and evolution. Astrochemistry labs are found in different environments ranging from dark clouds and protoplanetary disks to different phases of interstellar medium (ISM), and cluster formation in galaxies.

The astrophysical environment impacts its chemistry, and vice versa. For instance, cooling and heating processes as a result of chemical reactions will influence the dynamics. Hence, a hydrodynamics model needs to be coupled with a chemistry model, apart from

radiation, in order to fully simulate an astrophysical environment. Moreover, this dynamics is often complex and therefore requires a 3-dimensional approach when modelling it. 3D hydrodynamical modelling is a notoriously computationally expensive process, both in a particle-based and grid-based approach. It is often the case that coupling such hydrodynamics with a classical chemical model in every time step makes it computationally infeasible to explore even a modest physical parameter space. Various research groups have already made elaborate efforts in integrating (limited) chemistry in hydrodynamical simulations. To name a few, Glover & Mac Low (2007a,b); Walch et al. (2015), and Hu et al. (2021) combine hydro and chemistry, amongst other processes, in the case of molecular clouds and the ISM, Lahén et al. (2020) incorporated both constituents in star formation simulations, Yoneda et al. (2016) and Young et al. (2021) did so for protoplanetary disk research (the latter doing chemistry in a post-processing step), and Richings & Schaye (2016)

for galaxy formation,. In this research, we take an alternative pathway to produce feasible hydro-chemical simulations.

The astrophysical environment studied in this paper is the circumstellar envelope (CSE) of asymptotic giant branch (AGB) stars, i.e., evolved stars with an initial mass ranging between 0.8 and 8 solar masses. The CSE is created by the global stellar outflow launched at the surface of the AGB star (e.g., Freytag et al. 2017), believed to be due to the combination of surface pulsations facilitating dust formation (Bowen 1988; Höfner & Olofsson 2018). Mass-loss rates are typically found to be between 10^{-8} and $10^{-4} \text{ M}_{\odot} \text{ yr}^{-1}$, and terminal velocities are ranging from about 5 to 20 km s^{-1} (Knapp et al. 1998; Habing & Olofsson 2004; Ramstedt et al. 2009). Thanks to the favourable physical conditions and large physical gradients in the outflow, the CSE hosts a rich chemistry; over 100 molecules and about 15 dust species have been detected so far (Verhoelst et al. 2009; Waters 2011; Gail & Sedlmayr 2013; Decin 2021). Recently, CSEs have been found to contain asymmetrical structures, such as spirals, arcs, disks, and bipolarity (e.g., Maun & Huggins 2006; Kervella et al. 2016; Li et al. 2016; Decin 2021). The most probable hypothesis is that the gravitational interaction with an unseen (sub)stellar companion shapes the outflow, causing the asymmetrical morphology (Nordhaus & Blackman 2006; Decin et al. 2020; Gottlieb et al. 2022). 3D hydrodynamical simulations affirm this hypothesis (e.g. Theuns & Jorissen 1993; Mastrodemos & Morris 1999; Kim & Taam 2012; El Mellah et al. 2020; Maes et al. 2021; Malfait et al. 2021). While more complex physical processes are being implemented in hydrodynamical simulations step by step (Maes et al. 2022), such as radiation processes (e.g. Chen et al. 2017; Esseldeurs et al. 2023), no 3D hydro-chemical simulation of these environments exists yet, which is necessary in order to start bridging the gap between theory and observations.

Due to the expansion of CSEs, its chemical composition never reaches an equilibrium state. In these dynamical conditions, the evolution of the chemical abundances over time is described by a chemical kinetics model. More specifically, the change of number density of the chemical species in the desired chemical network is calculated by a set of non-linear, coupled ordinary differential equations (ODEs, e.g., Millar 2015). Chemical kinetics, from a mathematical point of view, is known to be a stiff problem (e.g., Wen et al. 2023), due to (i) the short time scales associated with certain chemical variation in astrophysical environments and hence the need for small time steps in the solver, and

(ii) the range in parameters that spans many orders of magnitude. Hence, computation is typically relatively slow for an extensive chemical network of a couple of hundred species.

The classical way to work around the long computation times is to reduce the extensive chemical network, given an astrophysical environment, to its most important species and reactions only, since typically only a handful of chemical species dominate the overall chemistry. In order to do so, different routines exist (e.g. Grassi et al. 2013). For CSEs of AGB stars, this has been done, for example, by Boulanger et al. (2019), for the chemistry in the inner part of the CSE with the aim to evaluate it simultaneously with hydrodynamics. In order to verify whether the reduced network properly represents the desired chemistry, one needs to perform an extensive parameter study over a large dataset and have error measures to quantify its performance. Reducing the chemical network has the advantage that it has a clear chemical interpretation: it leaves in the dominant chemical reactions that play a crucial role in the chemical kinetics, and removes reactions that do not contribute much. However, due to the non-linear nature of chemistry, secondary reactions, deemed unimportant at first sight, might be crucial in the end (e.g., as a result of chain reactions). Therefore, it is beneficial to be able to include an extended chemical network for thoroughly studying the chemical destruction and formation pathways of species in a given astrophysical environment.

Another way to speed up the computation time is by building a surrogate model, i.e., a model that is able to emulate the chemistry with a (severely) reduced computation time. In recent years, many different dimensionality reduction and function approximation techniques have been developed, commonly categorised within the field of Machine Learning (ML), that show promising prospects in the field of (astro)chemistry (Wen et al. 2023). In the framework of chemistry in the interstellar medium, Grassi et al. (2011) were the first to emulate chemical kinetics with a simple neural network (NN). Later, de Mijolla et al. (2019) have shown that an immense speed up, up to a factor of 10^5 , can be gained with this approach. Holdship et al. (2021) and Grassi et al. (2022) both used an autoencoder architecture (Kramer 1992) to first reduce the dimensionality of their chemical network and subsequently evolve the chemical kinetics on this reduced network. The way they evolve the chemistry however differs; while in Holdship et al. (2021) another NN handles this, Grassi et al. (2022) train the autoecoder to be interpretable still in the reduced chemical space, and solve the reduced chemistry

with a interpretable set of ODEs. [Sulzer & Buck \(2023\)](#) adopt a similar method as the latter, but step away from the physical interpretation of the reduced chemical network, and use a linear function to evolve the reduced features. Additionally, other ML strategies have been used to speed up the computation time of chemical models. [Branca & Pallottini \(2022\)](#) opted for physics informed neural networks (PINN, [Raissi et al. 2019](#)), and in a follow-up work ([Branca & Pallottini 2024](#)) they explored the use of a deep neural operator (DeepONet, [Lu et al. 2021](#)), while [Palud et al. \(2023\)](#) make use of NNs as a regression technique.

However, the crucial caveat in these works is that the time evolution of the chemistry is performed in a *static* physical environment. The objective of this work is to develop a scheme that is able to emulate the chemical evolution in a *dynamical* physical environment, namely with changing physical parameters over time.

As such, we introduce MACE – a *Machine learning Approach to Chemistry Emulation*. With MACE we step away from the classical way of calculating chemical evolution in a dynamic environment. We follow a similar approach as [Sulzer & Buck \(2023\)](#), and map the chemical kinetics problem to a reduced dynamical system that does no longer necessarily have a physical or chemical interpretation. Hence, we deliberately sacrifice the interpretability of our model to allow for more freedom to optimise the trade-off between the accuracy and the computational speed of the resulting surrogate model.

This paper is organised as follows. Sect. 2 describes the different methods to calculate chemical abundances. In Sect. 3, we introduce the architecture of our emulator MACE and in Sect. 4 its training methodology is presented in detail. Sect. 5 encompasses the results of different trained MACE models, which are further discussed in Sect. 6. In the latter we also suggest some future improvements. Finally, we conclude in Sect. 7.

2. METHODOLOGY

Before going into details about MACE, we first introduce the methodology of classical chemical kinetics, how this translates to emulation, and describe its application specifically on AGB circumstellar envelopes.

2.1. Classical chemical kinetics

Classically, given a physical state (commonly determined by density, temperature, and radiation field), the evolution over time of the abundances of the chemical species is described by a chemical kinetics model. This model describes the rate of change of the number density n by solving a set of coupled, non-linear

ordinary differential equations (ODEs). For a certain chemical species i , such a rate equation is given by

$$\frac{dn_i}{dt} = \sum_{j \in F_i} \left(k_j \prod_{r \in R_j} n_r \right) - \sum_{j \in D_i} \left(k_j \prod_{r \in R_j} n_r \right). \quad (1)$$

It states the balance between the formation reactions F_i (first term), and destruction reactions D_i (second term) of species i with reactant r from the set of reactants R_j , where k_j is the rate coefficient of reaction j . The rate $\frac{dn_i}{dt}$ has units $\text{cm}^{-3} \text{s}^{-1}$. If the chemical network only contains one-body and two-body reactions, the rate equation for species i reduces to

$$\frac{dn_i}{dt} = \sum_{j,l} k_{jl} n_j n_l + \sum_m k_m n_m - n_i \left(\sum_r k_{ir} n_r + \sum_s k_s \right), \quad (2)$$

where now the indices j, l, m, r, s run over all other chemical species in the network. We can write this more succinctly in matrix notation as

$$\frac{dn_\alpha}{dt} = A_{\alpha\beta} n_\beta + B_{\alpha\beta\gamma} n_\beta n_\gamma, \quad (3)$$

where we sum over repeated indices. $A_{\alpha\beta}$ is a matrix containing the one-body reaction coefficients k_m and k_s from Eq. (2), in units of s^{-1} , and $B_{\alpha\beta\gamma}$ is a tensor with the two-body reaction rates k_{jl} and k_{ir} as elements, in units of $\text{cm}^3 \text{s}^{-1}$. Henceforth, we will indicate matrices containing the coefficients in bold to clearly indicate their multi-dimensional nature. $\mathbf{A}$ and $\mathbf{B}$ are sparse, since chemical species only react with a limited number of other species.

2.2. Emulating chemical kinetics

We build a surrogate model that is able to emulate the chemical evolution given by Eq. (3). More specifically, we map the chemical kinetics system to a reduced dynamical system that can be solved more efficiently. The emulator takes the initial set of abundances $\mathbf{n}_0$ as input, together with the physical input $\mathbf{p}$. Subsequently, the emulator predicts the evolution of the chemical species over a given time t , i.e., the set of abundances $\hat{\mathbf{n}}(t)$, where the $\hat{\cdot}$ -symbol is used for predictions made by the emulator. In Sect. 3, we elaborate on its architecture.

2.3. Application: AGB circumstellar envelope

A 1-dimensional representation of an AGB star's circumstellar envelope is used as the astrochemical environment for the development of MACE. Since this environment is already well studied and understood (e.g., [Millar et al. 2000](#); [Li et al. 2016](#); [Van de Sande](#)

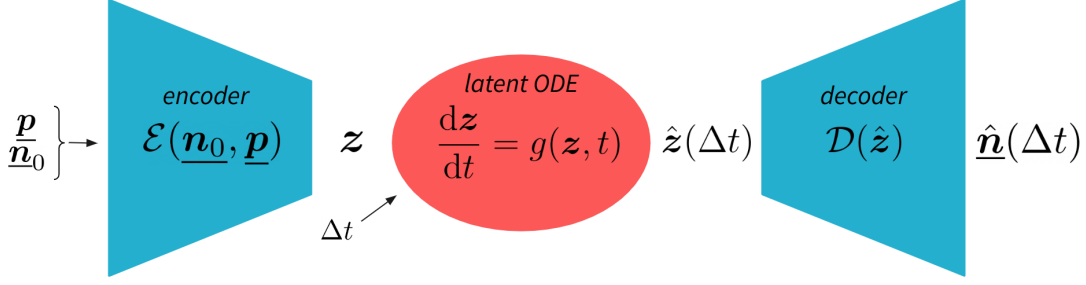


Figure 1: Schematic representation of the architecture of MACE. The autoencoder (encoder $\mathcal{E}$ and decoder $\mathcal{D}$) is given in blue. The latent ODE g is given in red. The flow of the surrogate model is from left to right. The initial abundances $\underline{n}_0$ and physical parameters $\underline{p}$ are concatenated and fed into the encoder, resulting in the latent representation $\underline{z}$. This latent representation is fed into the latent ODE function g , together with a timestep Δt . Solving this differential equation results in the predicted evolution in latent space $\hat{\underline{z}}(\Delta t)$. This predicted latent vector is then fed into the decoder, resulting in the predicted abundances $\hat{\underline{n}}(\Delta t)$.

et al. 2019; Maes et al. 2023), MACE can be benchmarked against this well-studied case. To generate training data, we use a version of the publicly available CSE model of the UMIST Database for Astrochemistry (UDfA)¹, which calculates the chemical abundances as a function of distance from the star (using a modified form of Eq. 2). We use an adaptation of the chemical network RATE12². The network consists of gas-phase chemistry only, involving 468 different chemical species (including electrons as a separate species), that can interact via 6180 reactions. Different types of reactions are included, such as two-body reactions between neutral and ionised species, photodissociation, and cosmic-ray induced reactions (for details, see McElroy et al. 2013). We will further refer to this model as the ‘1D-CSE’ model.

The 1D-CSE model assumes a power law for the density ρ and the temperature T , as a function of radius:

$$\rho(r) = \frac{\dot{M}}{4\pi r^2 v_{\text{exp}}}, \quad (4)$$

$$T(r) = T_{\star} \left(\frac{r}{R_{\star}} \right)^{-\varepsilon}. \quad (5)$$

The mass-loss rate $\dot{M}$ and the expansion velocity v_{exp} set the density, whereas the exponent ε sets the steepness of the temperature gradient. Here, $T_{\star}$ is the temperature of the AGB star at its surface $R_{\star}$. H_2 is assumed to be fully self-shielding. Further details about the model can be found in, e.g., Millar et al. (2000), Cordiner & Millar

(2009), McElroy et al. (2013), and Van de Sande et al. (2018). In Sect. 4.1, we elaborate on the specific dataset used for the training.

3. EMULATOR ARCHITECTURE

The architecture of the MACE emulator is chosen with two goals in mind: (i) we aim to take into account the full chemical network of RATE12, as such we want to reduce its dimensionality, (ii) we aim for a flexible emulator that is not restricted to a specific evolution time step. That is, the emulator should be able to accurately predict the evolution pathway of the chemical species over time without errors growing too much throughout the evolution. Hence, the architecture of MACE consists of two main parts: (i) an autoencoder for the dimensionality reduction, and (ii) a trainable ODE as substitute for the chemical evolution, on both we elaborate later in this section. Schematically, we can write MACE as the following function:

$$\hat{\underline{n}}(t) = \mathcal{D}\left(G(\mathcal{E}(\underline{n}, \underline{p}), t)\right), \quad (6)$$

where the underlined symbols indicate that that parameter is preprocessed and not used in its “raw” form (see further in Eq. 14 in Sect. 4.1). $\mathcal{E}$ and $\mathcal{D}$, the encoder and decoder respectively, constitute the autoencoder, while the function G describes the evolution in latent space such that $\underline{z}(\Delta t) = G(\underline{z}, \Delta t) = \int_0^{\Delta t} g(\underline{z}) dt$, with $\underline{z}$ the latent space variables. A schematic visualisation of the architecture is given in Fig. 1. MACE is built in Python using the PyTorch framework (Paszke et al. 2019) and is publicly available on GitHub: <https://github.com/silkemaes/MACE>.

3.1. Autoencoder

Autoencoders are a widely used tool to reduce the dimensionality of a certain set of data and typically

¹ The model can be found on GitHub: https://github.com/MarieVdS/rate22_cse_code (Millar et al. 2024)

² The network can be found online: <http://udfa.ajmarkwick.net/index.php?mode=downloads>.

Table 1: Number of nodes in the different layers of the autoencoder (encoder + decoder). The parameter d indicates the number of dimension in latent space.

	input	hidden 1	hidden 2	output
encoder	$468 + 4$	256	64	d
decoder	d	64	256	468

Notes. For all layers, except the output layer of the encoder, a leaky ReLU function with slope 0.2 is used as activation function. For the output layer of the encoder, a tanh is used.

consist of an encoder and a decoder, $\mathcal{E}$ and $\mathcal{D}$ in Eq. (6), respectively. They are a type of neural network architecture that are trained to reproduce their input as output. The encoder takes the input and maps it to a latent space, the lower-dimensional representation of the high-dimensional input. The decoder then maps the latent space back to the original input space. The latent space is thus a compressed representation of the input (Kramer 1992).

We construct a purely mathematical representation of the chemical network in latent space. Since in a chemical reaction network of 468 chemical species only a couple of a dozen of species are dominating the chemical pathways, it is expected that the dimensionality of the chemical network can be reduced by an order of magnitude at least, without losing important information (Holdship et al. 2021; Grassi et al. 2022; Sulzer & Buck 2023).

For MACE, the encoder $\mathcal{E}$ takes as input the set of chemical abundances $\underline{n}$, concatenated with the physical parameters $\underline{p}$, resulting in $(468 + 4)$ nodes in the input layer. Mathematically, it can be written as

$$\underline{z} = \mathcal{E}(\underline{n}, \underline{p}), \quad (7)$$

where $\underline{z}$ represents the parameters in latent space. The input layer is followed by two hidden layers and an output layer. The output layer has a number of nodes equal to the dimensionality of the latent space d , which is varied in the training stage (Sect. 4.3). The number of nodes in each hidden layer decreases and is chosen to equal a power of 2 in order to optimise computational resources, namely 256 and 64. The activation function in every layer, except for the output layer, is a leaky rectifier (leaky ReLU) with slope 0.2 (Maas 2013). For the output layer, a hyperbolic tangent (tanh) is used, in order to map the input to values within a range $]-1, 1[$ in latent space³.

³ We want to restrict the range in values in the latent space to better control the latent dynamics (see Sect. 3.3).

The decoder $\mathcal{D}$ has the same architecture as the encoder, but with the layers of the encoder reversed. In the decoder only leaky ReLUs with a slope of 0.2 are used as activation function. The number of nodes in the output layer of the decoder matches the number of chemical species. Mathematically, the decoder is given by

$$\hat{\underline{n}} = \mathcal{D}(\underline{z}). \quad (8)$$

To train the autoencoder, we aim for the following expression to be satisfied:

$$\underline{n} = \mathcal{D}(\mathcal{E}(\underline{n}, \underline{p})). \quad (9)$$

An overview of the layers and number of nodes in the autoencoder can be found in Table 1.

3.2. Latent ODE

Inspired by the mathematical form of the chemical rate equations, Eq. (2), we opted for the latent set of ODEs to be of the same form and include a constant term, akin to Sulzer & Buck (2023). The latent ODE is given by

$$\begin{aligned} \frac{dz_\alpha}{dt} &= g(z_\alpha, t) \\ &= \mathcal{C}_\alpha + \mathcal{A}_{\alpha\beta} z_\beta + \mathcal{B}_{\alpha\beta\gamma} z_\beta z_\gamma, \end{aligned} \quad (10)$$

where $\underline{z}$ is the encoded latent representation of $(\underline{n}, \underline{p})$, Eq. (7), and $\mathcal{C}$, $\mathcal{A}$, and $\mathcal{B}$ are constant tensors with dimensions matching the operation. The elements in these tensors are trainable parameters that are optimised during training. We solve the latent ODE (Eq. 10) using the PyTorch library *torchode* (Lienen & Günnemann 2022), allowing to train it at the same time due to its gradient tracking. Explicitly solving the latent ODE adds complex dynamics to the emulator, contrary to Sulzer & Buck (2023)’s approach, which makes it better grasp the complexity of the chemical kinetics model.

3.3. Latent dynamics

The dynamics of a general set of coupled, non-linear ODEs, such as the one given by Eq. (10), can be very chaotic without any constraints on the values of the coefficients. As such, the overall behaviour of the system is difficult to control. For instance, the evolution of a system as in Eq. (10), initialised with random coefficients, will typically make one or more components diverge. This will complicate the training process; on the one hand we want to explore the space of latent dynamics as much as possible, but on the other hand, the typical divergence of the system will regularly

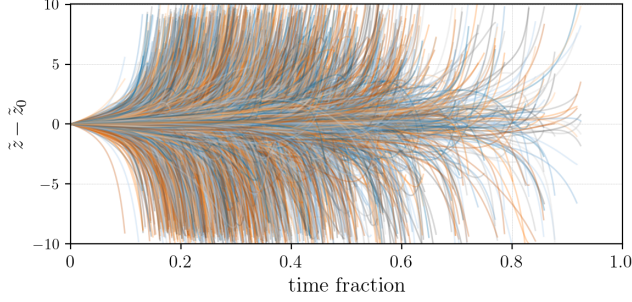


Figure 2: Example of result of an empirical latent dynamics test. The evolution as a function of time (given in time fraction) is given for 1000 randomly generated $\tilde{z}_8$, following Eq. (10).

cause problems in any attempt to numerically solve it. One way to resolve this is by constraining the latent dynamics only to “well-behaved” systems of ODEs that do not cause any divergences over the relevant time span. Although there is plenty of mathematical literature on the behaviour of dynamical systems that could help constrain the coefficients of our latent dynamics (e.g., [Strogatz 2000](#)), we opted for a more empirical approach.

Our approach is based on a simple observation: for any dynamical system of the form Eq. (10) with random (but finite) coefficients and initial conditions $\mathbf{z}_d(0)$ of dimension d , the width of the distribution of evolved latent variables $\mathbf{z}_d(t)$ can be controlled by the time t . In latent space, the time variable t no longer has any physical meaning, so we can rescale it to control the latent dynamics. The appropriate rescaling of the time variable will depend on the distributions of the ODE coefficients and the initial conditions, but most importantly on the dimension d of the dynamical system.

Determining the optimal rescaling is done with an empirical strategy using a numerical experiment. We produce 10 000 random initial values for the latent space vector, $\tilde{\mathbf{z}}_d(0)$ (the $\tilde{\cdot}$ -symbol indicating the empirical test), uniformly distributed between -1 and 1 to match the possible outcome of the encoder (see Sect. 3.1), as a function of different latent dimensionality d . Given the tensors $\tilde{\mathbf{C}}$, $\tilde{\mathbf{A}}$, and $\tilde{\mathbf{B}}$ with random, standard normally distributed values, Eq. (10) is solved for each $\tilde{\mathbf{z}}_d$ using *torchode* ([Lienen & Günnemann 2022](#)). An example of the outcome of such an empirical test for $d = 8$, i.e., the evolution of $\tilde{\mathbf{z}}_8$ over a normalised time, is shown in Fig. 2.

In order to determine the appropriate rescaling of the time parameter, the following methodology is used: (i) We require that a cutoff of 95% of empirical tests should still have a bound solution after a time t_{cutoff} . This

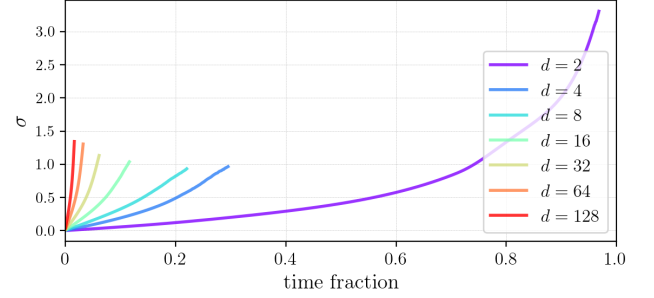


Figure 3: Standard deviations σ for the spread in $\tilde{\mathbf{z}}_d$ for the empirical tests of the latent dynamics (Eq. 10) for different dimensionality d , as a function of time fraction. The cutoff of each curve happens when, at that fraction in time, less than 95% of the empirical solutions are no longer bound and thus diverge to infinity.

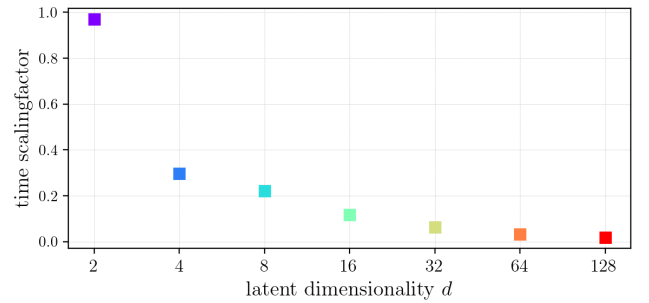


Figure 4: Resulting time rescaling per latent dimensionality d from empirical latent dynamics test. The colours correspond to the colours in Fig. 3.

time t_{cutoff} , normalised to the full evolution time, will serve as the scaling factor. (ii) We perform empirical tests for dimensionality $d \in [2, 4, 8, 16, 32, 64, 128]$. (iii) We look for the time t_{cutoff} by calculating the standard deviation σ as a function of evolution time. This σ serves as a measure of the range of the dynamics of an empirical latent system. The σ 's are only computed when at least 95% of the models still has a bound solution, giving us t_{cutoff} . Fig. 3 shows the σ 's for different latent dimensionality. The resulting time scale factors, after normalising the t_{cutoff} are shown in Fig. 4. For increasing dimensionality, the scaling fraction decreases, since a higher latent dimensionality increases the chance of a divergence.

4. TRAINING

Training an emulator is a non-trivial optimisation problem involving numerous free parameters. In this section, we elaborate on the parameter space, the data used for training, the different loss functions that are involved in the optimisation, and distinguish two training schemes.

4.1. Parameter space & dataset

To be efficient with data and computational resources, we reuse the grid of chemical models from the sensitivity analysis performed by [Maes et al. \(2023\)](#) as training data for MACE. The grid was constructed to span a broad range of CSE parameters found of AGB stars, based on observations. These 1D-CSE models were generated using the CSE model introduced in Sect. 2.3 and consists of 18 000 models of the chemistry in carbon-rich CSEs for varying densities and temperatures. The set of initial abundances (i.e., parent species) specific to carbon-rich AGB outflows can be found in Table 2.

Each of the models spans an outflow radius from 10^{14} to 10^{18} cm measured from the centre of the star, with a radial resolution of 134 steps. The ranges of the different parameters that were used to build the grid of 1D-CSE models are given in Table 3. A visualisation of the density and temperature space can be found

Table 2: Parent species of the carbon-rich AGB outflows, and their initial abundances relative to H_2 .

Species	Abundance
He	0.17
CO	8.00×10^{-4}
C_2H_2	4.38×10^{-5}
HCN	4.09×10^{-5}
N_2	4.00×10^{-5}
SiC_2	1.87×10^{-5}
CS	1.06×10^{-5}
SiS	5.98×10^{-6}
SiO	5.02×10^{-6}
CH_4	3.50×10^{-6}
H_2O	2.55×10^{-6}
HCl	3.25×10^{-7}
C_2H_4	6.85×10^{-8}
NH_3	6.00×10^{-8}
HCP	2.50×10^{-8}
HF	1.70×10^{-8}
H_2S	4.00×10^{-9}

Notes. Abundances taken from [Agúndez et al. \(2020\)](#). When a range was given there, the linear average is used.

Table 3: Physical parameters and their ranges of the grid of chemical models from [Maes et al. \(2023\)](#). The density is determined by $\dot{M}$ and v_{exp} according to Eq. (4) for the combinations given in Fig. 14, the temperature profile is given by the combination of $T_\star$ and ε through Eq. (5). The stellar radius, $R_\star$, inner radius, R_{inner} , and outer radius, R_{outer} , are kept constant. For the purpose of this research, the radii are converted to time, using the expansion velocity of the corresponding simulations. Adapted from [Maes et al. \(2023\)](#).

Parameter	Range/Value	Stepsize
$\dot{M}$	$[\text{M}_\odot \text{ yr}^{-1}] \quad 1 \times 10^{-8} - 5 \times 10^{-5}$	(*)
v_{exp}	$[\text{km s}^{-1}] \quad 2.5 - 25$	2.5
$T_\star$	$[\text{K}] \quad 2000 - 3000$	50
ε	/ $0.3 - 1.0$	0.05
$R_\star$	$[\text{cm}] \quad 2 \times 10^{13}$	/
R_{inner}	$[\text{cm}] \quad 10^{14}$	/
R_{outer}	$[\text{cm}] \quad 10^{18}$	/

Note. (*) For $\dot{M}$ we used the values 1×10^{-p} , 2×10^{-p} and 5×10^{-p} , with $p \in [5, 6, 7, 8]$, see also Fig. 14.

in Appendix A. Since the models assume a constant outflow velocity, we can relate the distance scale to time scale, using that $t = r/v_{\text{exp}}$. Hence, this results in a maximal chemical evolution time ranging from 0.4×10^{12} s to 4×10^{12} s for the range of expansion velocities considered.

The chemical evolution in the grid of models is determined mainly by four physical parameters

$$\mathbf{p} = (\rho, T, \xi, A_V), \quad (11)$$

where ρ is the density, T the temperature, and the two latter parameters denote the radiation field: ξ sets the mean interstellar radiation field ([Draine 1978](#)), normalised over a 3D sphere, and A_V is the outward dust extinction in the visible part of the spectrum, thus can be seen as an optical depth. As such, only the rate of photodissociation reactions is varied and not the cosmic-ray rate, since in smooth outflows it is found cosmic rays to not play a significant role (e.g., [Van de Sande & Millar 2019](#)). Hence, the vector $\mathbf{p}$ in Eq. (11) is the input for the physical parameters in the MACE architecture (Eq. 6).

To build the training dataset from the grid of chemistry models from [Maes et al. \(2023\)](#), the density ρ and temperature T are calculated as given by Eqs. (4) and (5), respectively. The two remaining parameters ξ and A_V depend indirectly on ρ . A_V is calculated as

$$A_V = \frac{A_{\text{UV}}}{[A_{\text{UV}}/A_V]}, \quad (12)$$

where A_{UV} is the extinction in the UV part of the electromagnetic spectrum and the ratio $[A_{\text{UV}}/A_{\text{V}}] = 4.65$ (Nejad et al. 1984). When we assume the extinction to be the same as that of the interstellar medium of $1.87 \times 10^{21} \text{ atoms cm}^{-2} \text{ mag}^{-1}$ (Cardelli et al. 1989), then A_{UV} is given by

$$A_{\text{UV}} = [A_{\text{UV}}/A_{\text{V}}] \frac{N_{\text{H}_2}}{1.87 \times 10^{21}}, \quad (13)$$

where $N_{\text{H}_2} = \rho_{\text{H}_2} \times r$ is the column density of H_2 . The parameter ξ is implemented as described by Jura & Morris (1981), calculating its value by taking the mean over a 3D sphere and depends on A_{UV} (Eq. 13). Fig. 16 in Appendix A shows the physical parameters $\mathbf{p}$ as a function of radius and time for an example 1D-CSE model, with $\dot{M} = 1 \times 10^{-6} \text{ M}_{\odot} \text{ yr}^{-1}$, $v_{\text{exp}} = 15 \text{ km s}^{-1}$, $T_{\star} = 2500 \text{ K}$, and $\varepsilon = 0.6$, resembling an average AGB outflow.

The data at separate timesteps in one 1D-CSE model are considered as separate training samples. Hence, each training sample has a constant set of physical parameters $\mathbf{p}$. We will refer to this training data as ‘0D-CSE’ samples. Since every 1D-CSE model has 134 steps in time, we have $18\,000 \times 134 = 2.41 \times 10^6$ 0D-CSE samples in total available for training, validating, and testing MACE.

The physical parameters of the 0D-CSE samples, as well as the abundances, span orders of magnitude. Generally, this makes it more difficult to accurately train an architecture. Before we feed the samples to MACE, we perform the following transformation for the physical parameters to bring the input values closer together in terms of order of magnitude.

$$\underline{\mathbf{p}} \equiv \mathcal{F}(\log_{10}(\mathbf{p})), \quad (14)$$

where $\mathcal{F}$ is the min-max rescaling function given by

$$\mathcal{F}(\mathbf{x}) = \frac{\mathbf{x} - \min(\mathbf{x})}{\max(\mathbf{x}) - \min(\mathbf{x})}, \quad (15)$$

which normalises the input and brings its values in the range $[0, 1]$. For the input abundances, we first clip them to 10^{-20} , since this was the absolute tolerance of the ODE solver of the 1D-CSE models. We then perform the same transformation as for the physical parameters, given in Eq. (14).

4.2. Loss functions

In order to train MACE, we define three localised losses and one integrated loss to express the desired behaviour of MACE that we want to optimise for.

- The *local absolute loss* (ABS) is defined as

$$\mathbf{L}_{\text{ABS}} = |\underline{\mathbf{n}}(t) - \hat{\underline{\mathbf{n}}}(t)|, \quad (16)$$

where $\underline{\mathbf{n}}$ and $\hat{\underline{\mathbf{n}}}$ are the real and predicted abundance, both transformed according to Eq. (14), respectively. The aim of the absolute loss is to enforce MACE to match the predicted values of the abundances with the real ones. Since the abundances of the different species are rescaled to a range of $[0, 1]$, the absolute loss will treat all species with a more equal importance compared to the unscaled abundances.

- The *local gradient loss* (GRD) is calculated as

$$\mathbf{L}_{\text{GRD}} = \left| \frac{d\underline{\mathbf{n}}}{dt} - \frac{d\hat{\underline{\mathbf{n}}}}{dt} \right|, \quad (17)$$

which aims to enforce that the evolution of the abundances matches with the classical model, rather than only the instantaneous values.

- The *identity loss* (IDN) is defined as

$$\mathbf{L}_{\text{IDN}} = |\underline{\mathbf{n}} - \mathcal{D}(\mathcal{E}(\underline{\mathbf{n}}, \mathbf{p}))|, \quad (18)$$

with $\mathcal{E}$ the encoder and $\mathcal{D}$ the decoder, respectively. It is used to enforce that the autoencoder reproduces the input as its output, and thus minimises the loss of information during the compression of the real chemical space to the latent space. Note here that the identity loss only acts on the real abundances $\underline{\mathbf{n}}$ at a certain time (and their un-evolved, autoencoded state), and not on the predicted abundances $\hat{\underline{\mathbf{n}}}$.

- We define a *time-integrated absolute loss* (iABS), which enforces to match a chemical *pathway* instead of only the next step in time. More specifically, this loss compares the predicted path (i.e., evolution of the chemical species over time) with the real chemical path. The integrated absolute loss is given by

$$\mathbf{L}_{\text{iABS}} = \sqrt{\int_t \left(\underline{\mathbf{n}}(t') dt' - \hat{\underline{\mathbf{n}}}(t') dt' \right)^2}, \quad (19)$$

where the integral goes over a time interval $[1, m]$. If $m = 1$, the integrated absolute loss reduces to the local absolute loss given in Eq. (16).

For each loss, we take the norm during training (since they are vectors containing the loss per chemical species) and sum over the different samples to obtain a mean squared error (MSE) per loss type:

$$L_{\text{type}} = \frac{1}{N} \sum_{i=1}^N (\mathbf{L}_{\text{type}, i})^2, \quad (20)$$

where the index i goes over the different samples in the dataset.

To optimise the free parameters in MACE, we follow two different schemes and compare in Sect. 5 which of the two schemes provides the proper method to predict the evolution of chemical abundances in a dynamical environment.

1. *Local training scheme*: training MACE on local information only. In this scheme, the total loss is defined as the sum of the local absolute loss (Eq. 16), the local gradient loss (Eq. 17), and the identity loss (Eq. 18), each weighted with a factor, below given by λ , as follows:

$$L_{\text{tot}}^{\text{local}} = \lambda_{\text{ABS}} L_{\text{ABS}} + \text{GRD} L_{\text{GRD}} + \lambda_{\text{IDN}} L_{\text{IDN}}. \quad (21)$$

The weights λ allow to boost certain loss types, in order to enforce certain behaviour of the model. After every timestep during the training, the free parameters of the model are updated. This method is visualised in the upper panel of Fig. 5: during the training, this loss tries to match the coloured data points (predicted) to the black curve (real), according to Eq. (21).

2. *Integrated training scheme*: training MACE using the integrated absolute loss (Eq. 19), hence taking into account the chemical evolution pathways. Since we still aim for the autoencoder to accurately encode and decode the data, also the identity loss (Eq. 18) is included during the integrated training:

$$L_{\text{tot}}^{\text{integrated}} = L_{\text{ABS}} + L_{\text{IDN}}, \quad (22)$$

where the losses are not weighted. The free parameters of the model are only updated after m timesteps during the training. The middle and bottom panels of Fig. 5 visualise this training scheme for $m = 16$ and $m = 64$, respectively. In this scheme, the predicted chemical path of length m is produced at every timestep, starting from the real abundance, using the MACE architecture (coloured curves). These predictions appear as “hairs” on top of the real abundance profile (black curves). The integrated absolute loss tries to match these predicted paths (the “hairs”) to the real chemical path (black curves), according to Eq. (22).

4.3. Hyperparameters

The chosen architecture results in a large number of hyperparameters, i.e., parameters specific to the

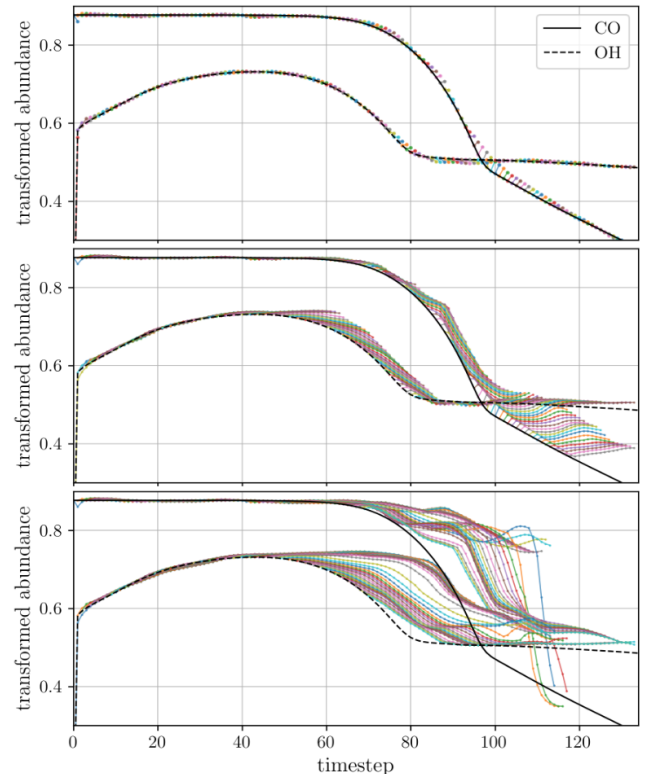


Figure 5: Visualisation of the integrated training scheme for different amount of timesteps m ; Upper panel: $m = 1$ (hence, this reduces to the local training). Middle panel: $m = 16$, Bottom panel: $m = 64$. During the training, the predicted chemical paths (coloured curves) are matched to the real path (black curves). An example is given for the abundance profiles of CO (full curve) and OH (dashed curve). The 1D-CSE model used here has the following parameters: $\dot{M} = 1 \times 10^{-8} \text{ M}_{\odot} \text{ yr}^{-1}$, $v_{\text{exp}} = 10 \text{ km s}^{-1}$, $T_{\star} = 2600 \text{ K}$, and $\varepsilon = 0.7$.

MACE architecture and training. Optimising for hyperparameters is a trial and error process, there is in principle no common way to do this. Moreover, it is possibly a degenerate problem, where changing different hyperparameters can have the same effect on the training and its results. Tools exist to do this in a systematic way. However, since this work serves a proof-of-concept, fully optimising the hyperparameters is beyond the scope of this work.

Therefore, we chose to keep a number of hyperparameters fixed, namely the number of layers and nodes in the autoencoder (see Sect. 3.1 and Table 1) and the learning rate. We test a limited amount of values for others, namely the dimensionality d of the latent space in both schemes, the weights of the losses

λ in the local training scheme (Eq. 21), and the number of timesteps m in the integrated training scheme.

4.4. Training strategy

The Adam optimiser (Kingma & Ba 2017), a stochastic gradient descent method, is used to train MACE. It is a popular choice for training neural networks, since it is computationally efficient, has little memory requirements, and is well suited for problems with large datasets and/or parameters. Moreover, it is an adaptive optimiser that adapts its stepsize in a particular way during gradient descent.

The different MACE models are trained for multiple epochs, where one epoch is defined as one full pass of the training and validation dataset. We chose to train for 100 epochs in the local training scheme (Eq. 21), and for 150 epochs in the integrated training scheme (Eq. 22), since the latter is more complex and is expected to need more epochs for the model to converge. At the start of the training, the losses included in the total loss differ by orders of magnitude. Therefore, after the first 5 epochs, we rescale them to unity in order to not favour a specific loss and skew the training. Subsequently, in the local training scheme, we scale each loss according to its weight λ .

In the local training scheme, we train four MACE architectures, each with a different latent dimensionality d and loss weights λ . The hyperparameters of the local MACE model, named *loc*, are given in Table 4. The integrated training scheme introduces an extra hyperparameter, namely the number of timesteps m of the chemical pathway. A large enough number of timesteps is preferred to enforce stability over a long enough period of time. However, introducing more timesteps will increase the computational cost of the training. We chose to train for three values: $m \in [8, 16, 32]$. Also, three different latent dimensionalities are used: $d \in [8, 16, 32]$. A larger latent dimensionality will result in more free parameters to train⁴, also increasing the computational cost of the training. Hence, by combining these values for m and d , we strive to find a balance between training time and accuracy. The combination of the different hyperparameters results in a grid of nine integrated MACE models, named *int*, given in Table 5.

The MACE models are trained on a random subset

⁴

For both the local and integrated models, when latent dimensionality $d = 8$, the number of free parameters equals 284 692, for $d = 16$ the number of free parameters is 289 508, and for $d = 32$ the number of free parameters is 321 028, giving an increase of about 36 000 between the former and the latter.

Table 4: Hyperparameters for the local MACE models, with d the dimensionality of the latent space, and λ the weights of the losses, as defined in Eq. (21). These models are trained for 100 epochs with an initial learning rate of 10^{-4} .

Model name	d	λ_{ABS}	λ_{GRD}	λ_{IDN}
<i>loc1</i>	8	1	1	1
<i>loc2</i>	8	10^4	10^2	10^2
<i>loc3</i>	16	1	1	1
<i>loc4</i>	16	10^4	10^2	10^2

Table 5: Hyperparameters for the integrated MACE models, with m the number of timesteps and d the dimensionality of the latent space. All integrated models are trained for 150 epochs with an initial learning rate of 10^{-4} .

Model name	m	d
<i>int1</i>	8	8
<i>int2</i>	8	16
<i>int3</i>	8	32
<i>int4</i>	16	8
<i>int5</i>	16	16
<i>int6</i>	16	32
<i>int7</i>	32	8
<i>int8</i>	32	16
<i>int9</i>	32	32

of 10 000 1D-CSE models (out of the 18 000), hence 1.34×10^6 0D-CSE samples, with a split of 70%–30% for training and validation data, respectively.

5. RESULTS

In this section, we show the results of trained MACE models with different hyperparameters and training strategies. We distinguish between the local (Sect. 5.1) and integrated training schemes (Sect. 5.2), each defined by their respective total loss function, given in Eqs. (21) and (22).

In order to test the trained MACE models, we apply them to a test dataset, containing 3 000 1D-CSE models, hence $3\,000 \times 134 \approx 4 \times 10^5$ 0D-CSE samples, separate from the training and validation set. We chose to apply the following metric on the test dataset for comparing the results from different models:

$$\text{error} = \frac{\log_{10} \mathbf{n} - \log_{10} \hat{\mathbf{n}}}{\log_{10} \mathbf{n}}, \quad (23)$$

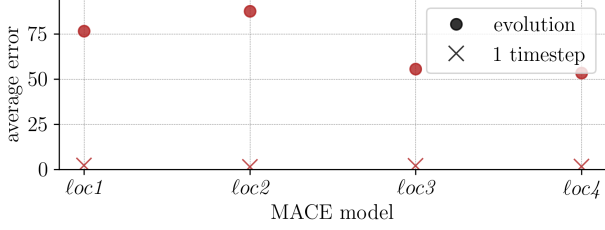


Figure 6: Average error (Eq. 23) on 4×10^5 test samples for the four local MACE models (Table 4). The dots give the error for testing the prediction of the full chemical evolution, the crosses give the error for testing the prediction of 1 consecutive timestep only.

which is executed element-wise and subsequently summed over the different chemical species. The metric states the relative “distance” in log-space between the true abundance $\mathbf{n}$ and the predicted abundance $\hat{\mathbf{n}}$.

5.1. Local training scheme

In this section, we discuss the four MACE models trained according to the local scheme (Eq. 21), given in Table 4. Fig. 6 gives the values of the error metric defined in Eq. (23), averaged, when applying the MACE models to the test dataset of 4×10^5 0D-CSE samples for predicting the full chemical evolution path (dots), and for predicting only the abundances at the consecutive timestep (crosses). Overall for the evolution prediction, this error is found to be rather large for the four different models. This indicates that these MACE models do not improve significantly when certain loss types are made to dominate. When the latent dimensionality is increased, the error is slightly lower. From the average errors (Fig. 6) we would prefer models *loc3* and *loc4* over the others, although not so very compelling. Fig. 7 shows the losses per epoch for model *loc3* (upper panel) and *loc4* (bottom panel). For model *loc3* it is seen that the identity loss dominates the training, which will also be the case for *loc1* (Table 4). For model *loc4* the weight of the absolute loss compared to the gradient and identity loss is increased, making the absolute loss dominate, as will be for *loc2*. Note that in Fig. 7 the losses have stabilised by training epoch 100, indicating a minimum has been reached, i.e., the optimal solution for these architectures. Note also that because the losses are rescaled before the training by a different factor for each models, their values have no absolute meaning, thus cannot be compared between models.

In order to verify if a certain MACE model is acceptable, apart from looking only at the value of the average error on test samples, we also consider its predicted abundances. As an example, we use the 1D-

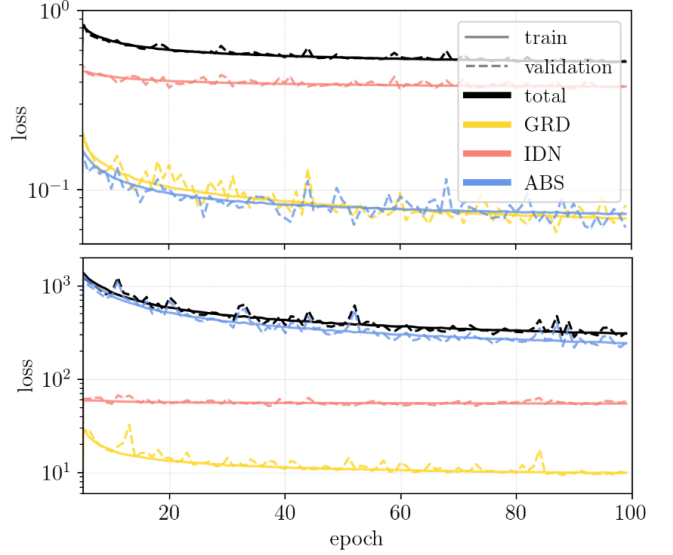


Figure 7: Losses per training epoch for model *loc3* (upper panel) and *loc4* (bottom panel). In colour are the individual losses (as indicated in the legend, abbreviated in Sect. 4.2). The black line gives the total loss, as defined in Eq. (21). The loss on the training data is given in full lines, the loss on the validation data in dashed.

CSE model with parameters $\dot{M} = 1 \times 10^{-6} M_{\odot} \text{ yr}^{-1}$, $v_{\text{exp}} = 17.5 \text{ km s}^{-1}$, $T_{\star} = 2300 \text{ K}$, and $\varepsilon = 0.55$ from the set of test models, resembling an average AGB outflow. To verify the evolution, we first feed the transformed input parameters of this 1D-CSE model, namely $\mathbf{p}_0$ and $\mathbf{n}_0$, together with the timestep $\Delta t_0 = t_1 - t_0$, to the trained MACE model and let it predict the next chemical state $\hat{\mathbf{n}}_1$ at t_1 . Subsequently, the MACE model reuses its chemical state $\hat{\mathbf{n}}_1$ and we feed it the corresponding physical state $\mathbf{p}_1$ from the classical model together with Δt_1 . The MACE model then predicts the successive chemical state $\hat{\mathbf{n}}_2$ at t_2 , and so on. Hence, the successive application of MACE predicts the chemical evolution for a dynamical physical environment.

As an example, we elaborate on the results of model *loc3*. Fig. 8 shows the predicted chemical evolution in full lines, and in dashed the classical analogue (taken here as “ground truth”). The middle panel shows the abundances, relative to H_2 , of three parent species (CO , H_2O , and C_2H_2), two first-generation daughter species (OH and C_2H), and two later-generation daughter species ($\text{CH}_3\text{C}_5\text{NH}^+$ and $\text{C}_{10}\text{H}_2^+$). The choice of displayed species is rather random within each category, however, the MACE architecture will not distinguish. The bottom panel shows the error metric, defined in Eq. (23), per species. Although MACE works with time as a parameter, in these figures time is converted back to

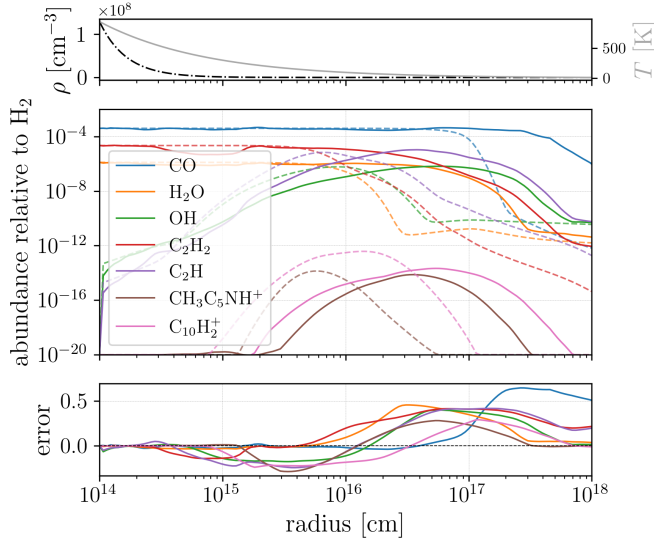


Figure 8: Evolution test of model *loc3* on a 1D-CSE models with the following input parameters: $\dot{M} = 1 \times 10^{-6} \text{ M}_{\odot} \text{ yr}^{-1}$, $v_{\text{exp}} = 17.5 \text{ km s}^{-1}$, $T_{\star} = 2300 \text{ K}$, and $\varepsilon = 0.55$. Upper panel: H_2 number density (dashed-dotted, left y -axis) and temperature (full grey, right y -axis) as a function of outflow radius. Middle panel: Abundances of specific species, given in legend. The dashed line gives the result of the classical model (ground truth), the full line gives the result of MACE. Lower panel: Error (Eq. 23) of the MACE model compared to the classical model.

outflow radius. We see that model *loc3* is able to predict the abundances of the parent species quite well until the abundance starts to drop. This is not very surprising, because up until that point, the model should not change the initial value much, since “doing nothing” is easy. However, further down the evolutionary path, we see that the MACE model diverges strongly from the ground truth. For the daughter species, this MACE model is not able to reproduce their chemical path. Also the other local models are not able to reproduce of the chemical pathway, figures are given in Appendix B.

Although the local MACE models fail to reproduce the chemical evolution, they do accurately reproduce 1 consecutive chemical state (1 timestep only). We test this by feeding the MACE model the physical state $\mathbf{p}_i$ and *true* chemical state $\mathbf{n}_i$ of the corresponding 1D-CSE model at every timestep t_i , and let it predict the next chemical state $\hat{\mathbf{n}}_{i+1}$ at t_{i+1} . An example of such a test is shown in Fig. 9 for model *loc3*, where the abundance predictions (dots) are shown in consecutive order according to the radius of the CSE. Within the chosen metric (Eq. 23), the error here is an order of magnitude lower than when predicting the full chemical

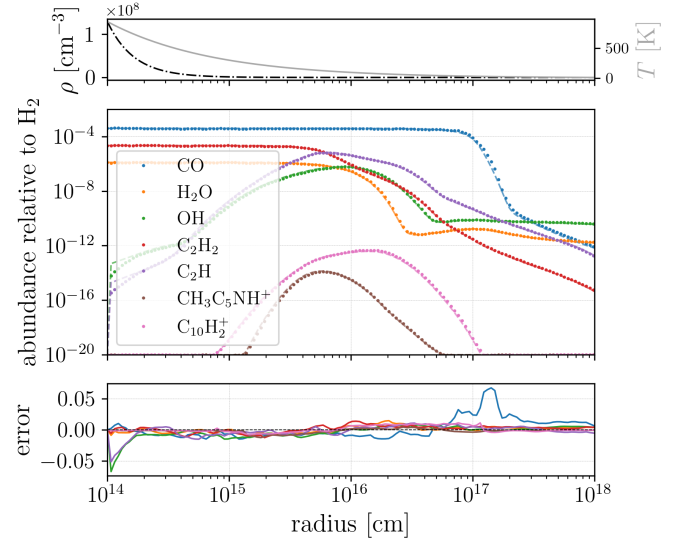


Figure 9: Timestep test of model *loc3* on a 1D-CSE models with the following input parameters: $\dot{M} = 1 \times 10^{-6} \text{ M}_{\odot} \text{ yr}^{-1}$, $v_{\text{exp}} = 17.5 \text{ km s}^{-1}$, $T_{\star} = 2300 \text{ K}$, and $\varepsilon = 0.55$. Upper panel: H_2 number density (dashed-dotted, left y -axis) and temperature (full grey, right y -axis) as a function of outflow radius.. Middle panel: Abundances of specific species, given in legend. The dashed line gives the result of the classical model (ground truth), the dots gives the result of MACE only used on 1 timestep, and shown in consecutive order according to the physical parameters of the CSE. Lower panel: Error (Eq. 23) of the MACE model compared to the classical model.

evolution, as is the case for the average error over the test dataset for the different local models (see crosses in Fig. 6). This is an impressive results, since, to the best of our knowledge, it is the first time that the next chemical step for such a large chemical network can be accurately predicted in a machine learning approach. However, this leaves us with the challenge to robustly predict the chemical evolution in a dynamical environment, which we address in the next section.

5.2. Integrated training scheme

In this section, we discuss the MACE models trained according to the integrated scheme (Eq. 22), given in Table 5. Fig. 10 shows the error metric (Eq. 23), averaged over the same test dataset of 4×10^5 samples, after applying the MACE models on it. Within this metric, the average errors are about a factor 4 smaller than the average errors on the local models (Fig. 6). This is a significant improvement. Again, the errors of the different models lie close to each other, from which we cannot clearly prefer one over the other. Fig.

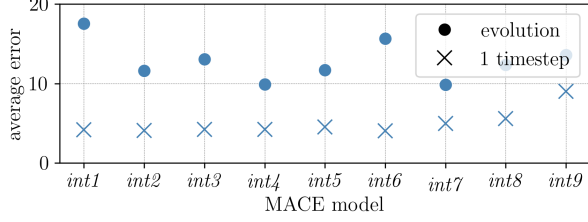


Figure 10: Average error (Eq. 23) on 4×10^5 test samples for the integrated MACE models (Table 5). The dots give the error for testing the prediction of the full chemical evolution, the crossing give the error for testing the prediction of 1 consecutive timestep only.

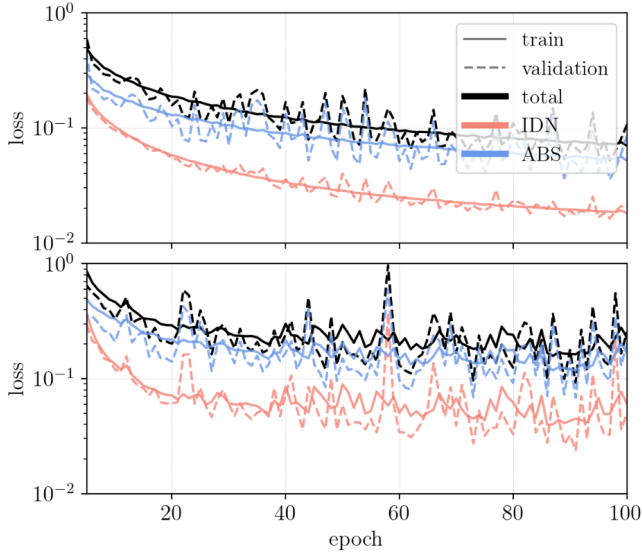


Figure 11: Losses per training epoch for model *int2* (upper panel) and *int7* (bottom panel). In colour are the individual losses (as indicated in the legend, abbreviated in Sect. 4.2). The black line gives the total loss, as defined in Eq. (21). The loss on the training data is given in full lines, the loss on the validation data in dashed.

11 shows examples of the loss per epoch for models *int2* (upper panel) and *int7* (bottom panel). We see that for model *int7* the training loss (full lines) does not decrease with every epoch, giving it a spiky look. This model is more complex, containing many more free parameters than, e.g., *int2* (see footnote 4). Therefore, it is harder to find a local minimum, especially when the learning rate is large. We suspect that training this model with a lower learning rate will smoothen the train loss curve. Though, the general trend of the training loss is a decrease, indicating that the model does converge.

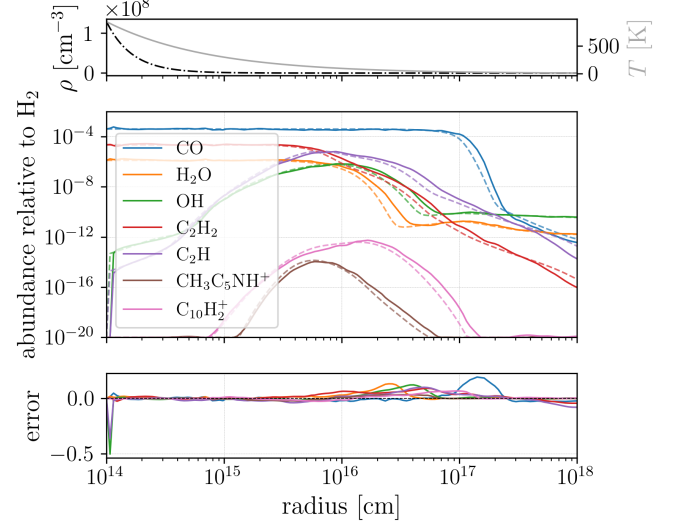


Figure 12: Evolution test of model *int4* on a 1D-CSE models with the following input parameters: $\dot{M} = 1 \times 10^{-6} M_{\odot} \text{ yr}^{-1}$, $v_{\text{exp}} = 17.5 \text{ km s}^{-1}$, $T_{\star} = 2300 \text{ K}$, and $\varepsilon = 0.55$. Upper panel: H_2 number density (dashed-dotted, left y -axis) and temperature (full grey, right y -axis) as a function of outflow radius. Middle panel: Abundances of specific species, given in legend. The dashed line gives the result of the classical model (ground truth), the full line gives the result of MACE. Lower panel: Error (Eq. 23) of the MACE model compared to the classical model.

Fig. 12 shows the predicted chemical evolution by model *int4* on the same example 1D-CSE test model. Compared to the results of the local models (Sect. 5.1), these findings are a great improvement. The predicted chemical evolution almost matches the real evolution exactly; the MACE model is able to reproduce the chemical pathway of the parent and daughter species. The other integrated models give similar, good results. Figures can be found in Appendix C.

Generally, the integrated MACE models perform better on a high density outflow. For lower density models, a systematic offset is noticed in the predicted abundances by MACE (see e.g., left panel of Fig. 19 in Appendix C). This is due to the effect of the CO self-shielding, which is depending on the velocity of the outflow and can affect the abundances significantly (Maes et al. 2023). However, since a MACE model receives the density as input (Eq. 11) and not the mass-loss rate and expansion velocity separately (contrary to the classical model), this chemical subtlety is not grasped by MACE. Though, this can easily be included in an improved version of MACE, as demonstrated by this proof-of-concept.

6. DISCUSSION

In this section, we elaborate on the integrated MACE models, since the integrated training scheme is found to be succesful, contrary to the local scheme. We discuss the accuracy of their results on test data, as a function of their training time. Moreover, we determine the speed-up of using MACE instead of the classical CSE model from Sect. 2.3. Finally, we discuss possible improvements of the MACE architecture and training.

6.1. Accuracy & training time

Training a more complex model generally takes more time. As such, it is only beneficial to train a more complex model when it results in a significant improvement in accuracy. Fig. 13 shows the performance landscape of our models; the mean error (Eq. (23), used here as a measure of accuracy) of the integrated MACE models is given as a function of the time it took to train one epoch⁵. The arrows indicate the direction of improvement within this landscape. The more complex models (i.e., models containing more free parameters and/or using more steps m in their training scheme) are indicated with a darker shade of blue. We find that they do not necessarily perform better than the simpler models (e.g., model *int9* versus model *int3*). More specifically, this indicates that using more timesteps m in the integrated training scheme is not necessary. Moreover, we find that the models with a smaller latent dimensionality d (e.g., *int2*, *int4*, and *int7*) give a better performance, as they have a better accuracy for a shorter training time.

Besides, there is a caveat regarding the performance of chemistry emulators. They are trained on data coming from classical chemistry models, in this case believed to be the “ground truth”. However, also the classical models give a certain error compared to reality, since they rely on reaction rates gathered from experiments or estimated from theory. Assessing the error on classical models via sensitivity analysis is an active field of research (e.g., Van de Sande, *in prep*) and ballpark values are yet unknown. Hence, there is no use in trying to reduce the mean error on MACE to zero. Accordingly, we claim *int2* to be our best performing model in this proof-of-concept work, since it is located in the performance landscape amongst the lowest mean errors and, especially, has the shortest training time.

⁵ Training times of the different models should be compared only relatively, since, depending on which and how many CPUs the training is performed, the absolute values will differ from what is given here. The models in this work are all trained on the same machine, using the same amount of CPUs

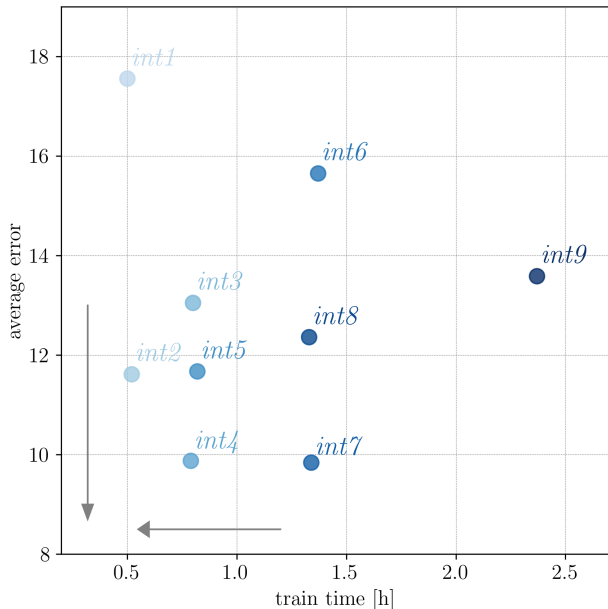


Figure 13: Performance landscape of integrated MACE models: accuracy (given by mean error, Eq. 23) as a function of training time needed to complete one epoch. The arrows indicate the direction of improvement in this landscape of the preferred models. The darker the colour, the more complex the model.

Table 6: Computation time to produce a 1D-CSE chemistry model with MACE and the classical DVODE solver. The computation time is given in seconds.

computation time [s]	MACE	classical	speed-up
low density	0.6	14.6	$\times 24$
high density	0.6	16.8	$\times 28$

6.2. Speed-up

The ultimate goal of building a chemical emulator is to speed up the chemistry computations. In this section, we elaborate on the speed-up that MACE offers compared to the classical CSE model. We time how long it takes the solvers (MACE versus the classical model, the latter using DVODE from ODEPACK Brown et al. 1989; Hindmarsh 1983) to calculate the evolution of the abundances of a certain 1D-CSE model, and as such exclude any overhead time. We distinguish between a 1D-CSE model with a low and high density, as we know the computation of the DVODE solver will differ in both cases. The results are given in Table 6.

With MACE, we find a speed-up of a factor 24 for low-density outflows and a factor 28 for high-

density outflows, since the MACE solver is independent of the density in the outflow. Although a speed-up of one order of magnitude for chemistry calculations might not be sufficient to make a 3D hydro-chemistry feasible, the power of MACE lies within the optimised use of matrix operations in its implementation with PyTorch. As such, we can use the PyTorch framework to efficiently parallelise the computations for multiple particles. Furthermore, we can leverage hardware accelerators, such as graphics processing units (GPUs) or tensor processing units (TPUs), specifically optimised for this kind of computations. Therefore, the speed-up reported here provides merely a lower bound. The true practical speed up will depend on the particular application.

6.3. Future prospects

MACE serves as a proof-of-concept for emulating chemical evolution pathways in a dynamical environment. As such, its architecture can be trained on chemical models of other astrophysical environments as well, for instance protoplanetary disks, diffuse or molecular ISM, dark clouds, etc. The set of physical parameters used in this research (Eq. 11), most probably, will not be ideal to use in other environments. Therefore, the physical parameters should be adapted to the specific environment, a flexibility which MACE is designed for (Eq. 6).

Moreover, the current implementation of MACE can be further developed before it is applied in a hydrodynamical simulation. In this section, we discuss some of the possible improvements that can be made to the architecture and training in order to get a better performance, and elaborate on the application of MACE in hydrodynamical simulations.

6.3.1. Possible improvements

The architecture of MACE can be improved in different ways, from a physical/chemical point-of-view, as well as from an implementation/machine learning point-of-view.

Improving on the physics and chemistry in MACE can be done by including other types of (physically informed) losses in the training. For example, for an environment undergoing chemical evolution, the total mass of the environment is conserved (in the absence of nucleosynthesis). Grassi et al. (2022) and Sulzer & Buck (2023) have included mass conservation as an extra loss term in their architecture. This requirement can even be made stronger, by stating that the total number of atoms of each element is conserved, thus adding element conservation. Including this can be done by adding an extra loss term to the total loss (Eq. 21)

or by constructing the latent ODE (Eq. 10) in such a way that it is built in. In Appendix D, we elaborate on this prospect.

We can also improve MACE from a machine learning point-of-view. First of all, keeping the current architecture, the hyperparameters can be further optimised, for example by training with another learning rate and for more epochs. Secondly, the architecture can be developed further by adding more layers and/or nodes per layer to the encoder and decoder, allowing for more complexity. However, this will also add time to the training procedure. Therefore, it would be advantageous to leverage GPU acceleration, instead of CPU only. Additionally, the autoencoder and latent ODE can be trained separately, contrary to what is done in this work. This approach can be very beneficial; (i) it would allow for analysing and finetuning better the latent space and as such optimising its dynamics, and (ii) multiple timesteps can be performed in latent space without the need to decode and encode every step. This will not only save some computation time, but rather allow for economic memory use, since only a greatly reduced amount of features need to be stored. Expanding on this approach, multiple, various decoders can be developed to match the desired purpose. For example, if the purpose is to only model a set of parent species, a decoder can be trained to only return that specific set, again saving computation time. We note that, in order to separately train the two parts, the physical parameters (Eq. 11) should be incorporated in the latent space and not in the encoder, altering slightly the flow of MACE given in Eq. (6). This can be done, for example, by constructing the latent ODE coefficient tensors (Eq. 10) from neural networks taking the physical parameters as input.

Moreover, improvements in the implementation can benefit and increase the speed-up MACE provides over classical models. Currently, *torchode* (Lienen & Günnemann 2022) is used as ODE solver, since we need its gradient tracking to train the latent ODE system (Eq. 10). However, the solver by itself is relatively slow and currently dominates the computation time of MACE over the autoencoder by a factor of 10. Hence, once an optimal MACE model is acquired, *torchode* can be substituted for an alternative ODE solver that is faster for this dynamics. Furthermore, the tolerances of the ODE solver can be optimised regarding the error between the results of the emulator and the classical model, potentially increasing even more the speed-up.

6.3.2. Implementation in hydrodynamical simulations

Now that we have established that MACE performs properly and is able to reproduce a 1D chemical pathway in dynamical environment, the next step is to couple the emulator to a 3D hydrodynamics simulation. Because the chemical models, on which MACE is currently trained, cannot deal with the advection of chemical abundances, MACE should be coupled to a Lagrangian hydrodynamics model, which allows us to solve the chemistry in a co-moving reference frame, for example a smoothed particle hydrodynamics (SPH) framework (Lucy 1977; Gingold & Monaghan 1977). Moreover, in order to correctly implement the radiation-induced chemical reactions in a 3D hydrodynamical model, such as photodissociation, the radiation-related parameters, ξ and A_V (Eq. 11), should be calculated as well. This can either be achieved in a very approximate way, or more elaborate by the use of a ray-tracing algorithm.

As future work, we aim to couple MACE to the SPH code PHANTOM (Price & Federrath 2010; Price et al. 2018) in the framework of AGB outflows (Siess et al. 2022; Esseldeurs et al. 2023). The radiation-related parameters, ξ and A_V , will be calculated by using the ray tracer of the 3D radiative transfer solver MAGRITTE⁶ (De Ceuster et al. 2020a,b, 2022), similar to how it was used by Esseldeurs et al. (2023). The coupling will allow to generate 3D hydro-chemical models of AGB outflow perturbed by a companion (Maes et al. 2022). This would allow us to step away from a 1D-approach (Van de Sande & Millar 2022) and better study the impact of the companion on the chemistry in the outflow. As such, chemical signatures of hidden companions can be identified, which is crucial for the interpretation of observations of AGB outflows. Besides, the MACE framework provides us with the prospect of implementing chemical cooling and heating processes in the hydrodynamics simulation, and inform molecular line cooling, in order to model and study the interplay between chemistry and hydrodynamics, and more correctly model the physical processes in the outflow.

7. CONCLUSION

This work presents MACE, a *Machine learning Approach to Chemistry Emulation*, as a proof-of-concept for emulating chemistry in a dynamical environment. Inspired by literature findings (e.g., Holdship et al. 2021; Grassi et al. 2022; Sulzer & Buck 2023), we

have constructed an architecture (Eq. 6) where an autoencoder compresses the chemical network to a latent space (Eqs. 7 and 8). Subsequently, in this mathematical latent space, the chemical evolution is emulated by solving the latent coupled ordinary differential equation (Eq. 10). The physical parameters of the environment (Eq. 11) are included in the encoder.

For the first time, it is possible to accurately predict chemical abundances for a large chemical network of 468 species and 6180 reactions. Moreover, we find that using an integrated training scheme (Eq. 22) allows to reproduce a full chemical pathway in a *dynamical* environment, something that has not been done before. As an example, we apply it to the dynamical environment of AGB star’s circumstellar envelopes, with the objective to couple this chemistry emulator to existing 3D hydrodynamical models. In order for this to be feasible, MACE should have a fast performance.

We find MACE to have a speed-up of one order of magnitude compared to its classical analogue, complementary to the optimised use of matrix operations in its implementation with PyTorch. The latter makes that only one pass of the emulator is needed for computations of the chemistry of all particles in a mesh-free hydrodynamics simulation, contrary to the a classical chemistry model where computation will grow linearly.

The current implementation of MACE offers opportunities for further development, such as including element conservation and refining its architecture, most likely enhancing its performance. Moreover, MACE is designed to be flexible, so that it can be applied to other astrophysical environments as well, by retraining its architecture on the appropriate models.

ACKNOWLEDGEMENTS

S.M. and L.D. acknowledge support from the Research Foundation Flanders (FWO) grant G099720N. F.D.C. is a Postdoctoral Research Fellow of the Research Foundation - Flanders (FWO), grant number 1253223N, and was previously supported for this research by a Postdoctoral Mandate (PDM) from KU Leuven, grant number PDMT2/21/066. M.V.d.S. acknowledges support from the European Union’s Horizon 2020 research and innovation programme under the Marie Skłodowska-Curie grant agreement No 882991 and the Oort Fellowship at Leiden Observatory. L.D. also acknowledges support from KU Leuven C1 MAESTRO grant C16/17/007, KU Leuven C1 BRAVE grant C16/23/009, and KU Leuven Methusalem grant METH24/012.

⁶ MAGRITTE is open source and can be found online: <https://github.com/Magritte-code/Magritte>, <https://magritte.readthedocs.io/en/stable/>.

Software: UMIST 1D-CSE model (https://github.com/MarieVdS/rate22_cse_code, McElroy et al. 2013;

Millar et al. 2024), PyTorch (Paszke et al. 2019), *torchode* (Lienen & Günnemann 2022)

APPENDIX

A. PHYSICAL PARAMETER SPACE

This section shows the physical parameter space of the classical 1D-CSE models. Fig. 14 shows the number density ranges ($\rho_N = \frac{\rho}{\mu m_H}$, with ρ given in Eq. (4), μ the mean molecular mass per H_2 molecule and m_H the atomic mass unit) for the grid of models set by mass-loss rate $\dot{M}$ and expansion velocity v_{exp} , at three different locations in the outflow ($r = 10^{14}$ cm, $r = 10^{16}$ cm, and $r = 10^{18}$ cm). Fig. 15 shows the temperature ranges, set by $T_\star$ and ε in Eq. (5). Fig. 16 gives an example of the physical parameters (Eq. 11) of an 1D-CSE models as a function of radius and time, which forms the basis of the dynamical physical environment in which the chemical reactions take place.

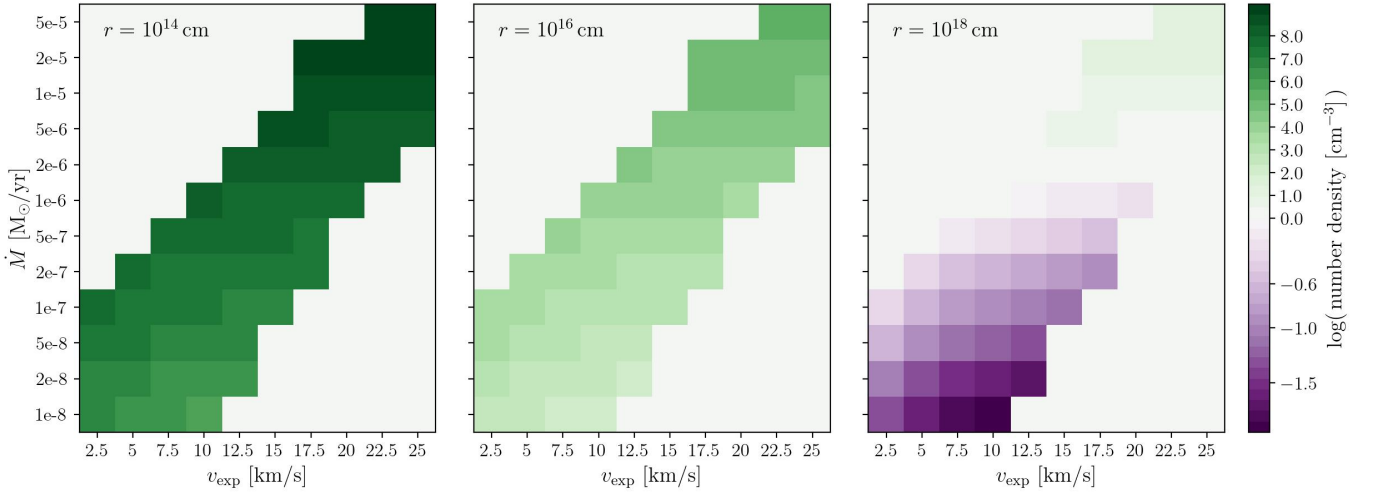


Figure 14: Visualisation of the number density space ($\rho_N = \frac{\rho}{\mu m_H}$, with ρ given in Eq. (4), μ the mean molecular mass per H_2 molecule and m_H the atomic mass unit) of the training data, via the combinations of expansion velocity v_{exp} and mass-loss rate $\dot{M}$, given at different radii. (Adapted from Maes et al. 2023)

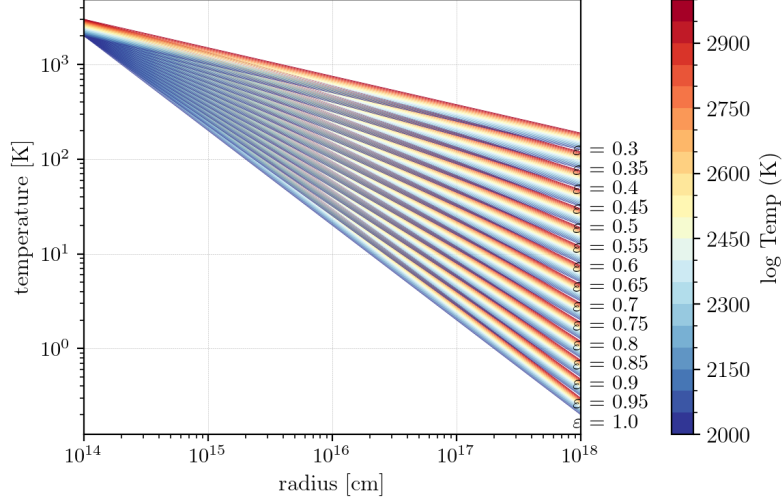


Figure 15: Visualisation of the different temperature profiles (Eq. 5) in the training data, where the stellar temperature T_* is indicated by the colourbar. The different values of ε result in different groups of temperature profiles, indicated at the right-hand side of the panel. (Adapted from Maes et al. 2023)

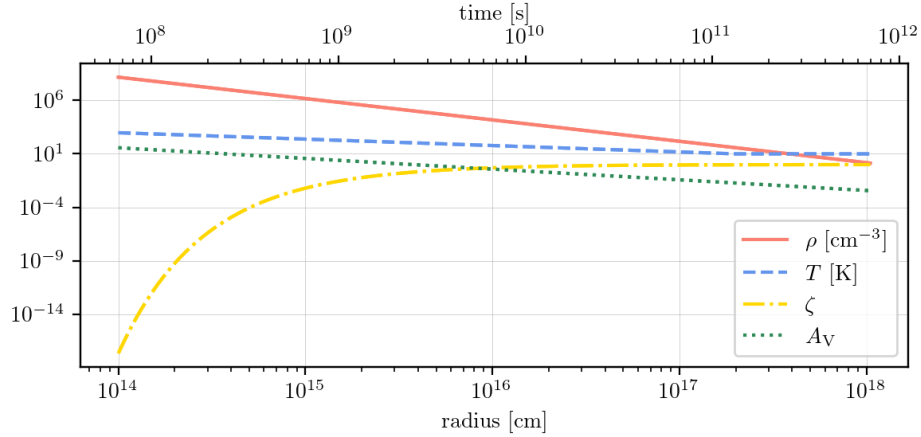


Figure 16: The physical parameters $\mathbf{p} = (\rho, T, \xi, A_V)$ as a function of time (upper x -axis) and radius (bottom x -axis), given for an example 1D-CSE models with input parameters $\dot{M} = 1 \times 10^{-6} \text{ M}_\odot \text{ yr}^{-1}$, $v_{\text{exp}} = 15 \text{ km s}^{-1}$, $T_* = 2500 \text{ K}$, and $\varepsilon = 0.6$. The y -axis states the value of the specific parameter given in by the legend.

B. TESTING LOCAL MODELS

In this section, we show more predicted abundance profiles by the MACE models, trained according to the local scheme (Eq. 21 and Table 4). Fig. 17 shows model *loc1* (left) and *loc4* (right) apply on a 1D-CSE test model.

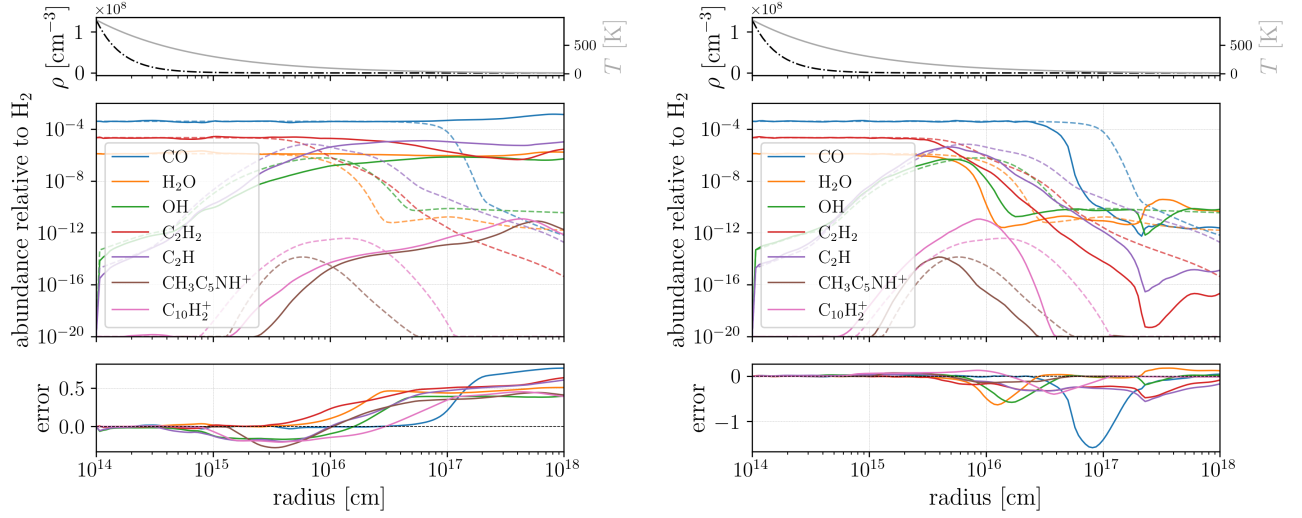


Figure 17: Evolution test of model *loc1* (left) and *loc4* (right) on a 1D-CSE models with the following input parameters: $\dot{M} = 1 \times 10^{-6} M_{\odot} \text{ yr}^{-1}$, $v_{\text{exp}} = 17.5 \text{ km s}^{-1}$, $T_{\star} = 2300 \text{ K}$, and $\varepsilon = 0.55$. Upper panel: H_2 number density (dashed-dotted, left y -axis) and temperature (full grey, right y -axis) as a function of outflow radius. Middle panel: Abundances of specific species, given in legend. The dashed line gives the result of the classical model (ground truth), the full line gives the result of MACE. Lower panel: Error (Eq. 23) of the MACE model compared to the classical model.

C. TESTING INTEGRATED MODELS

In this section, we show more predicted abundance profiles by the MACE models, trained using the integrated scheme (Eq. 22 and Table 5). Fig. 18 shows the predicted abundance profile by models *int2* (left) and *int7* (right). Fig. 19 shows tests of model *int4* on a low- (left) and high-density (right) outflow.

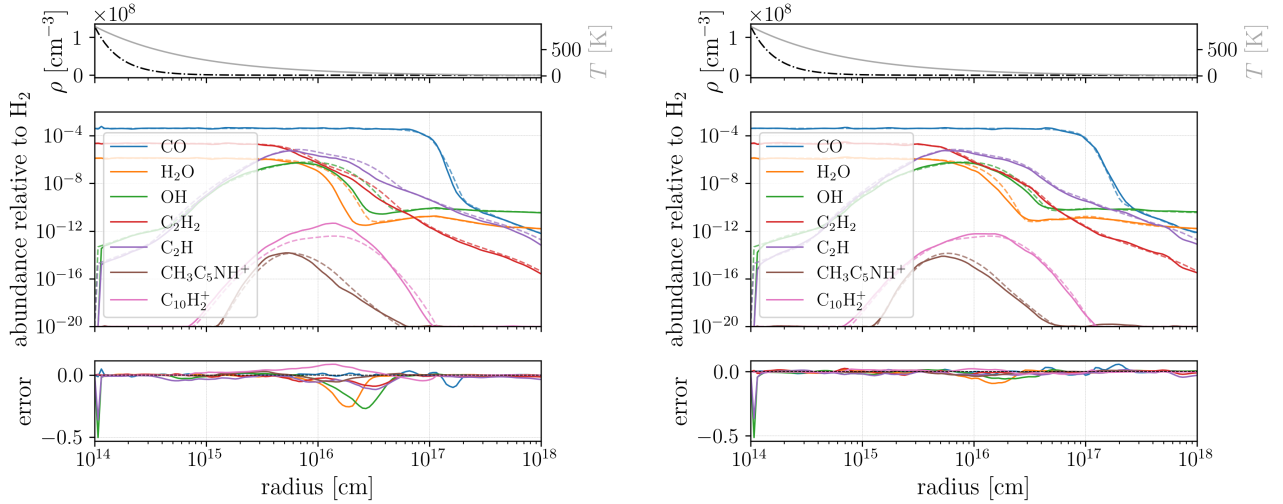


Figure 18: Evolution test of model *int2* (left) and *int7* (right) on a 1D-CSE models with the following input parameters: $\dot{M} = 1 \times 10^{-6} M_{\odot} \text{ yr}^{-1}$, $v_{\text{exp}} = 17.5 \text{ km s}^{-1}$, $T_{\star} = 2300 \text{ K}$, and $\varepsilon = 0.55$. Upper panel: H_2 number density (dashed-dotted, left y -axis) and temperature (full grey, right y -axis) as a function of outflow radius. Middle panel: Abundances of specific species, given in legend. The dashed line gives the result of the classical model (ground truth), the full line gives the result of MACE. Lower panel: Error (Eq. 23) of the MACE model compared to the classical model.

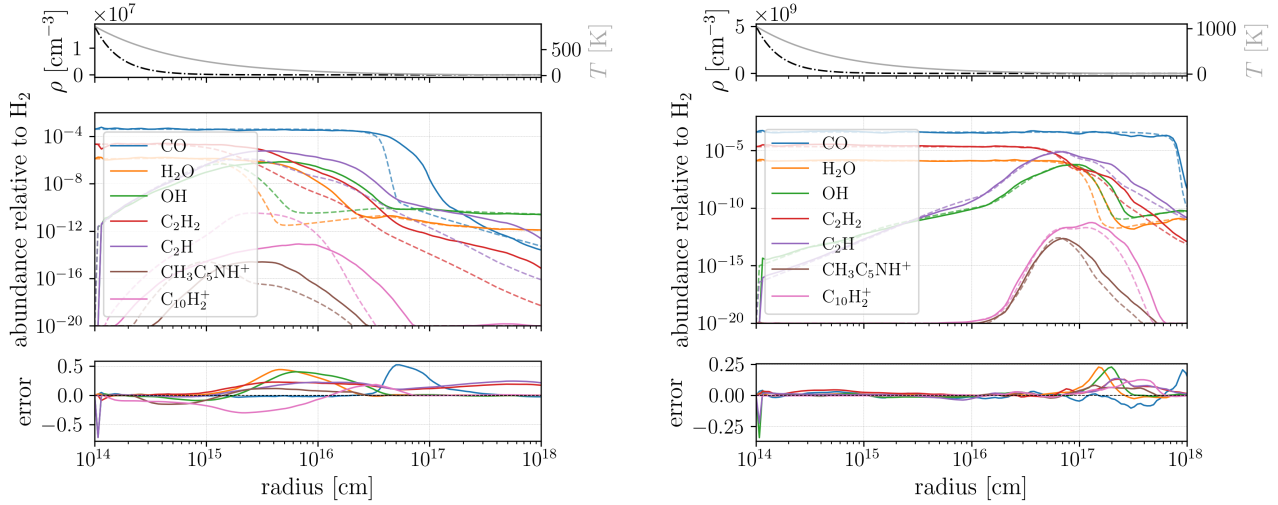


Figure 19: Evolution tests of model *int4*. Left: low-density 1D-CSE models with input parameters $\dot{M} = 2 \times 10^{-8} M_{\odot} \text{ yr}^{-1}$, $v_{\text{exp}} = 2.5 \text{ km s}^{-1}$, $T_{\star} = 2100 \text{ K}$, and $\varepsilon = 0.5$. Right: high-density 1D-CSE models with input parameters $\dot{M} = 5 \times 10^{-5} M_{\odot} \text{ yr}^{-1}$, $v_{\text{exp}} = 22.5 \text{ km s}^{-1}$, $T_{\star} = 2750 \text{ K}$, and $\varepsilon = 0.6$. Upper panel: H_2 number density (dashed-dotted, left y -axis) and temperature (full grey, right y -axis) as a function of outflow radius. Middle panels: Abundances of specific species, given in legend. The dashed line gives the result of the classical model (ground truth), the full line gives the result of MACE. Lower panels: Error (Eq. 23) of the MACE model compared to the classical model.

D. IMPLEMENTING ELEMENT CONSERVATION

In this section, we elaborate on the implementation of element conservation in the emulator. The matrix M_{Ii} defines how much of each element, I , appears in each chemical species, i , of the considered chemical network. The elemental abundance e_I then yields $M_{Ii} n_i$ and should be conserved, hence

$$M_{Ii} \frac{dn_i}{dt} = 0. \quad (\text{D1})$$

Given (i) that the decoder maps latent abundances to real abundances, $n_i = \mathcal{D}_i(\mathbf{z})$ (Eq. 8) and (ii) the dynamics in latent space (Eq. 10), we can rewrite the dynamics in real space as

$$\frac{dn_i}{dt} = \partial_{\alpha} n_i \frac{dz_{\alpha}}{dt} = \partial_{\alpha} \mathcal{D}_i(\mathbf{z}) \frac{dz_{\alpha}}{dt} = \partial_{\alpha} \mathcal{D}_i(\mathbf{z}) (\mathcal{C}_{\alpha} + \mathcal{A}_{\alpha\beta} z_{\beta} + \mathcal{B}_{\alpha\beta\gamma} z_{\beta} z_{\gamma}). \quad (\text{D2})$$

Using the dynamics in real space, conservation of elemental abundance (Eq. 10) thus implies

$$M_{Ii} \partial_{\alpha} \mathcal{D}_i(\mathbf{z}) (\mathcal{C}_{\alpha} + \mathcal{A}_{\alpha\beta} z_{\beta} + \mathcal{B}_{\alpha\beta\gamma} z_{\beta} z_{\gamma}) = 0, \quad (\text{D3})$$

with $\partial_{\alpha} = \partial/\partial z_{\alpha}$.

The conservation of elemental abundance can be implemented in the emulator in two ways. (i) The conservation of elements can be present as an additional loss term to the total loss. Hence, the element loss (ELM) would be defined as

$$L_{\text{ELM}} = M_{Ii} \partial_{\alpha} \mathcal{D}_i(\mathbf{z}) (\mathcal{C}_{\alpha} + \mathcal{A}_{\alpha\beta} z_{\beta} + \mathcal{B}_{\alpha\beta\gamma} z_{\beta} z_{\gamma}). \quad (\text{D4})$$

We have tried implementing this element loss in MACE. However, this slowed down the training immensely, because for every pass of training data through MACE, the jacobian of the decoder neural network $\partial \mathcal{D}(\mathbf{z})$ needs to be calculated and multiplied by the tensor coefficients of the latent ODE, which involves many operations on large matrices. (ii) The conservation of elements can be incorporated explicitly in the architecture itself, by constructing the latent ODE (Eq. 10) in such a way that Eq. (D3) is always satisfied. This can be done in many different ways, requiring many explicit design choices, such as which parameters to determine using this constraint, which is beyond the scope of this paper.

REFERENCES

- Agúndez, M., Martínez, J. I., de Andres, P. L., Cernicharo, J., & Martín-Gago, J. A. 2020, *A&A*, 637, A59, doi: [10.1051/0004-6361/202037496](https://doi.org/10.1051/0004-6361/202037496)
- Boulangier, J., Clementel, N., van Marle, A. J., Decin, L., & de Koter, A. 2019, *MNRAS*, 482, 5052, doi: [10.1093/mnras/sty2560](https://doi.org/10.1093/mnras/sty2560)
- Bowen, G. H. 1988, *ApJ*, 329, 299, doi: [10.1086/166378](https://doi.org/10.1086/166378)
- Branca, L., & Pallottini, A. 2022, *Monthly Notices of the Royal Astronomical Society*, 518, 5718, doi: [10.1093/mnras/stac3512](https://doi.org/10.1093/mnras/stac3512)
- Branca, L., & Pallottini, A. 2024, arXiv e-prints, arXiv:2402.12435, doi: [10.48550/arXiv.2402.12435](https://doi.org/10.48550/arXiv.2402.12435)
- Brown, P. N., Byrne, G. D., & Hindmarsh, A. C. 1989, *SIAM J. Sci. Stat. Comput.*, 10, 1038
- Cardelli, J. A., Clayton, G. C., & Mathis, J. S. 1989, *ApJ*, 345, 245, doi: [10.1086/167900](https://doi.org/10.1086/167900)
- Chen, Z., Frank, A., Blackman, E. G., Nordhaus, J., & Carroll-Nellenback, J. 2017, in *IAU Symposium*, Vol. 323, *Planetary Nebulae: Multi-Wavelength Probes of Stellar and Galactic Evolution*, ed. X. Liu, L. Stanghellini, & A. Karakas, 367–368, doi: [10.1017/S1743921317001715](https://doi.org/10.1017/S1743921317001715)
- Cordiner, M. A., & Millar, T. J. 2009, *ApJ*, 697, 68, doi: [10.1088/0004-637X/697/1/68](https://doi.org/10.1088/0004-637X/697/1/68)
- De Ceuster, F., Homan, W., Yates, J., et al. 2020a, *MNRAS*, 492, 1812, doi: [10.1093/mnras/stz3557](https://doi.org/10.1093/mnras/stz3557)
- De Ceuster, F., Bolte, J., Homan, W., et al. 2020b, *MNRAS*, 499, 5194, doi: [10.1093/mnras/staa3199](https://doi.org/10.1093/mnras/staa3199)
- De Ceuster, F., Ceulemans, T., Srivastava, A., et al. 2022, *Journal of Open Source Software*, 7, 3905, doi: [10.21105/joss.03905](https://doi.org/10.21105/joss.03905)
- de Mijolla, D., Viti, S., Holdship, J., Manolopoulou, I., & Yates, J. 2019, *A&A*, 630, A117, doi: [10.1051/0004-6361/201935973](https://doi.org/10.1051/0004-6361/201935973)
- Decin, L. 2021, *ARA&A*, 59, 337, doi: [10.1146/annurev-astro-090120-033712](https://doi.org/10.1146/annurev-astro-090120-033712)
- Decin, L., Montargès, M., Richards, A. M. S., et al. 2020, *Science*, 369, 1497, doi: [10.1126/science.abb1229](https://doi.org/10.1126/science.abb1229)
- Draine, B. T. 1978, *ApJS*, 36, 595, doi: [10.1086/190513](https://doi.org/10.1086/190513)
- El Mellah, I., Bolte, J., Decin, L., Homan, W., & Keppens, R. 2020, *A&A*, 637, A91, doi: [10.1051/0004-6361/202037492](https://doi.org/10.1051/0004-6361/202037492)
- Esseldeurs, M., Siess, L., De Ceuster, F., et al. 2023, *A&A*, 674, A122, doi: [10.1051/0004-6361/202346282](https://doi.org/10.1051/0004-6361/202346282)
- Freytag, B., Liljegren, S., & Höfner, S. 2017, *A&A*, 600, A137, doi: [10.1051/0004-6361/201629594](https://doi.org/10.1051/0004-6361/201629594)
- Gail, H.-P., & Sedlmayr, E. 2013, *Physics and Chemistry of Circumstellar Dust Shells*
- Gingold, R. A., & Monaghan, J. J. 1977, *MNRAS*, 181, 375, doi: [10.1093/mnras/181.3.375](https://doi.org/10.1093/mnras/181.3.375)
- Glover, S. C. O., & Mac Low, M.-M. 2007a, *ApJS*, 169, 239, doi: [10.1086/512238](https://doi.org/10.1086/512238)
- . 2007b, *ApJ*, 659, 1317, doi: [10.1086/512227](https://doi.org/10.1086/512227)
- Gottlieb, C. A., Decin, L., Richards, A. M. S., et al. 2022, *A&A*, 660, A94, doi: [10.1051/0004-6361/202140431](https://doi.org/10.1051/0004-6361/202140431)
- Grassi, T., Bovino, S., Schleicher, D., & Gianturco, F. A. 2013, *MNRAS*, 431, 1659, doi: [10.1093/mnras/stt284](https://doi.org/10.1093/mnras/stt284)
- Grassi, T., Merlin, E., Piovan, L., Buonomo, U., & Chiosi, C. 2011, arXiv e-prints, arXiv:1103.0509, doi: [10.48550/arXiv.1103.0509](https://doi.org/10.48550/arXiv.1103.0509)
- Grassi, T., Nauman, F., Ramsey, J. P., et al. 2022, *A&A*, 668, A139, doi: [10.1051/0004-6361/202039956](https://doi.org/10.1051/0004-6361/202039956)
- Habing, H. J., & Olofsson, H. 2004, *Asymptotic Giant Branch Stars*, doi: [10.1007/978-1-4757-3876-6](https://doi.org/10.1007/978-1-4757-3876-6)
- Hindmarsh, A. C. 1983, *Scientific Computing*, 1, 55
- Höfner, S., & Olofsson, H. 2018, *A&A Rv*, 26, 1, doi: [10.1007/s00159-017-0106-5](https://doi.org/10.1007/s00159-017-0106-5)
- Holdship, J., Viti, S., Haworth, T. J., & Ilee, J. D. 2021, *A&A*, 653, A76, doi: [10.1051/0004-6361/202140357](https://doi.org/10.1051/0004-6361/202140357)
- Hu, C.-Y., Sternberg, A., & van Dishoeck, E. F. 2021, *ApJ*, 920, 44, doi: [10.3847/1538-4357/ac0dbd](https://doi.org/10.3847/1538-4357/ac0dbd)
- Jura, M., & Morris, M. 1981, *ApJ*, 251, 181, doi: [10.1086/159452](https://doi.org/10.1086/159452)
- Kervella, P., Homan, W., Richards, A. M. S., et al. 2016, *A&A*, 596, A92, doi: [10.1051/0004-6361/201629877](https://doi.org/10.1051/0004-6361/201629877)
- Kim, H., & Taam, R. E. 2012, *ApJ*, 744, 136, doi: [10.1088/0004-637X/744/2/136](https://doi.org/10.1088/0004-637X/744/2/136)
- Kingma, D. P., & Ba, J. 2017, *Adam: A Method for Stochastic Optimization*. <https://arxiv.org/abs/1412.6980>
- Knapp, G. R., Young, K., Lee, E., & Jorissen, A. 1998, *ApJS*, 117, 209, doi: [10.1086/313111](https://doi.org/10.1086/313111)
- Kramer, M. 1992, *Computers & Chemical Engineering*, 16, 313, doi: [https://doi.org/10.1016/0098-1354\(92\)80051-A](https://doi.org/10.1016/0098-1354(92)80051-A)
- Lahén, N., Naab, T., Johansson, P. H., et al. 2020, *ApJ*, 891, 2, doi: [10.3847/1538-4357/ab7190](https://doi.org/10.3847/1538-4357/ab7190)
- Li, X., Millar, T. J., Heays, A. N., et al. 2016, *A&A*, 588, A4, doi: [10.1051/0004-6361/201525739](https://doi.org/10.1051/0004-6361/201525739)
- Lienen, M., & Günnemann, S. 2022, in *The Symbiosis of Deep Learning and Differential Equations II*, *NeurIPS*. <https://openreview.net/forum?id=uiKVKTiUYB0>
- Lu, L., Jin, P., Pang, G., Zhang, Z., & Karniadakis, G. E. 2021, *Nature Machine Intelligence*, 3, 218–229, doi: [10.1038/s42256-021-00302-5](https://doi.org/10.1038/s42256-021-00302-5)
- Lucy, L. B. 1977, *AJ*, 82, 1013, doi: [10.1086/112164](https://doi.org/10.1086/112164)
- Maas, A. L. 2013. <https://api.semanticscholar.org/CorpusID:16489696>
- Maes, S., Homan, W., Malfait, J., et al. 2021, *A&A*, 653, A25, doi: [10.1051/0004-6361/202140823](https://doi.org/10.1051/0004-6361/202140823)

- Maes, S., Van de Sande, M., Danilovich, T., De Ceuster, F., & Decin, L. 2023, MNRAS, 522, 4654, doi: [10.1093/mnras/stad1152](https://doi.org/10.1093/mnras/stad1152)
- Maes, S., Siess, L., Homan, W., et al. 2022, in *The Origin of Outflows in Evolved Stars*, ed. L. Decin, A. Zijlstra, & C. Gielen, Vol. 366, 227–233, doi: [10.1017/S1743921322000217](https://doi.org/10.1017/S1743921322000217)
- Malfait, J., Homan, W., Maes, S., et al. 2021, A&A, 652, A51, doi: [10.1051/0004-6361/202141161](https://doi.org/10.1051/0004-6361/202141161)
- Mastrodemos, N., & Morris, M. 1999, ApJ, 523, 357, doi: [10.1086/307717](https://doi.org/10.1086/307717)
- Mauron, N., & Huggins, P. J. 2006, A&A, 452, 257, doi: [10.1051/0004-6361:20054739](https://doi.org/10.1051/0004-6361:20054739)
- McElroy, D., Walsh, C., Markwick, A. J., et al. 2013, A&A, 550, A36, doi: [10.1051/0004-6361/201220465](https://doi.org/10.1051/0004-6361/201220465)
- Millar, T. J. 2015, *Plasma Sources Science and Technology*, 24, 043001, doi: [10.1088/0963-0252/24/4/043001](https://doi.org/10.1088/0963-0252/24/4/043001)
- Millar, T. J., Herbst, E., & Bettens, R. P. A. 2000, MNRAS, 316, 195, doi: [10.1046/j.1365-8711.2000.03560.x](https://doi.org/10.1046/j.1365-8711.2000.03560.x)
- Millar, T. J., Walsh, C., Van de Sande, M., & Markwick, A. J. 2024, A&A, 682, A109, doi: [10.1051/0004-6361/202346908](https://doi.org/10.1051/0004-6361/202346908)
- Nejad, L. A. M., Millar, T. J., & Freeman, A. 1984, A&A, 134, 129
- Nordhaus, J., & Blackman, E. G. 2006, MNRAS, 370, 2004, doi: [10.1111/j.1365-2966.2006.10625.x](https://doi.org/10.1111/j.1365-2966.2006.10625.x)
- Palud, P., Einig, L., Le Petit, F., et al. 2023, A&A, 678, A198, doi: [10.1051/0004-6361/202347074](https://doi.org/10.1051/0004-6361/202347074)
- Paszke, A., Gross, S., Massa, F., et al. 2019, arXiv e-prints, arXiv:1912.01703, doi: [10.48550/arXiv.1912.01703](https://doi.org/10.48550/arXiv.1912.01703)
- Price, D. J., & Federrath, C. 2010, MNRAS, 406, 1659, doi: [10.1111/j.1365-2966.2010.16810.x](https://doi.org/10.1111/j.1365-2966.2010.16810.x)
- Price, D. J., Wurster, J., Tricco, T. S., et al. 2018, PASA, 35, e031, doi: [10.1017/pasa.2018.25](https://doi.org/10.1017/pasa.2018.25)
- Raissi, M., Perdikaris, P., & Karniadakis, G. 2019, *Journal of Computational Physics*, 378, 686, doi: [https://doi.org/10.1016/j.jcp.2018.10.045](https://doi.org/https://doi.org/10.1016/j.jcp.2018.10.045)
- Ramstedt, S., Schöier, F. L., & Olofsson, H. 2009, A&A, 499, 515, doi: [10.1051/0004-6361/200911730](https://doi.org/10.1051/0004-6361/200911730)
- Richings, A. J., & Schaye, J. 2016, MNRAS, 458, 270, doi: [10.1093/mnras/stw327](https://doi.org/10.1093/mnras/stw327)
- Siess, L., Homan, W., Toupin, S., & Price, D. J. 2022, A&A, 667, A75, doi: [10.1051/0004-6361/202243540](https://doi.org/10.1051/0004-6361/202243540)
- Strogatz, S. H. 2000, *Nonlinear Dynamics and Chaos: With Applications to Physics, Biology, Chemistry and Engineering* (Westview Press)
- Sulzer, I., & Buck, T. 2023, arXiv e-prints, arXiv:2312.06015, doi: [10.48550/arXiv.2312.06015](https://doi.org/10.48550/arXiv.2312.06015)
- Theuns, T., & Jorissen, A. 1993, MNRAS, 265, 946, doi: [10.1093/mnras/265.4.946](https://doi.org/10.1093/mnras/265.4.946)
- Van de Sande, M., & Millar, T. J. 2019, ApJ, 873, 36, doi: [10.3847/1538-4357/ab03d4](https://doi.org/10.3847/1538-4357/ab03d4)
- . 2022, MNRAS, 510, 1204, doi: [10.1093/mnras/stab3282](https://doi.org/10.1093/mnras/stab3282)
- Van de Sande, M., Sundqvist, J. O., Millar, T. J., et al. 2018, A&A, 616, A106, doi: [10.1051/0004-6361/201732276](https://doi.org/10.1051/0004-6361/201732276)
- Van de Sande, M., Walsh, C., Mangan, T. P., & Decin, L. 2019, MNRAS, 490, 2023, doi: [10.1093/mnras/stz2702](https://doi.org/10.1093/mnras/stz2702)
- Verhoelst, T., van der Zypen, N., Hony, S., et al. 2009, A&A, 498, 127, doi: [10.1051/0004-6361/20079063](https://doi.org/10.1051/0004-6361/20079063)
- Walch, S., Girichidis, P., Naab, T., et al. 2015, MNRAS, 454, 238, doi: [10.1093/mnras/stv1975](https://doi.org/10.1093/mnras/stv1975)
- Waters, L. B. F. M. 2011, in *Astronomical Society of the Pacific Conference Series*, Vol. 445, *Why Galaxies Care about AGB Stars II: Shining Examples and Common Inhabitants*, ed. F. Kerschbaum, T. Lebzelter, & R. F. Wing, 227
- Wen, M., Spotte-Smith, E. W. C., Blau, S. M., et al. 2023, *Nature Computational Science*, 3, 12, doi: [10.1038/s43588-022-00369-z](https://doi.org/10.1038/s43588-022-00369-z)
- Yoneda, H., Tsukamoto, Y., Furuya, K., & Aikawa, Y. 2016, ApJ, 833, 105, doi: [10.3847/1538-4357/833/1/105](https://doi.org/10.3847/1538-4357/833/1/105)
- Young, A. K., Alexander, R., Walsh, C., et al. 2021, MNRAS, 505, 4821, doi: [10.1093/mnras/stab1675](https://doi.org/10.1093/mnras/stab1675)