

Bayesian Methods for the Investigation of Temperature Dependence in Conductivity

Andrew R. McCluskey^{1,2,*}, Samuel W. Coles^{3,4} and Benjamin J. Morgan^{5,6,†}

¹Centre for Computational Chemistry, School of Chemistry, *University of Bristol*, Cantock's Close, Bristol, BS8 1TS, United Kingdom

²*Diamond Light Source*, Harwell Campus, Didcot, OX11 0DE, United Kingdom

³Yusuf Hamied Department of Chemistry, *University of Cambridge*, Lensfield Road, Cambridge, CB2 1EW, United Kingdom

⁴Lennard-Jones Centre, *University of Cambridge*, Trinity Lane, Cambridge, CB2 1TN, United Kingdom

⁵Department of Chemistry, *University of Bath*, Claverton Down, Bath, BA2 7AY, United Kingdom

⁶*The Faraday Institution, Quad One*, Harwell Science and Innovation Campus, Didcot, OX11 0RA, United Kingdom



(Received 19 December 2025; published 21 May 2026)

Temperature-dependent transport data, including diffusion coefficients and ionic conductivities, are routinely analyzed by fitting empirical models such as the Arrhenius equation. These fitted models yield parameters such as the activation energy, and can be used to extrapolate to temperatures outside the measured range. Researchers frequently face challenges in this analysis: quantifying the uncertainty of fitted parameters, assessing whether the data quality is sufficient to support a particular empirical model, and using these models to predict behavior at temperatures outside the measured range. Bayesian methods offer a coherent framework that addresses all of these challenges. This tutorial introduces the use of Bayesian methods for analyzing temperature-dependent transport data, covering parameter estimation, model selection, and extrapolation with uncertainty propagation, with illustrative examples from molecular dynamics simulations of superionic materials.

DOI: [10.1103/nxgq-lmp6](https://doi.org/10.1103/nxgq-lmp6)

CONTENTS

I. INTRODUCTION	1
II. CHALLENGES IN TEMPERATURE-DEPENDENT MODELLING	2
A. Model selection	2
B. Parameter estimation	2
C. Extrapolation	3
III. PARAMETER ESTIMATION	3
IV. MODEL SELECTION	5
V. EXTRAPOLATION	6
VI. CONCLUSIONS	7
VII. METHODS	7
ACKNOWLEDGMENTS	8
DATA AVAILABILITY	8

APPENDIX A: MEAN-SQUARED DISPLACEMENT OF LLZO AS A FUNCTION OF TEMPERATURE	8
APPENDIX B: TEMPERATURE DEPENDENCE AS A FUNCTION OF DIFFUSIVE SIMULATION TIME FOR AgCrSe ₂	9
REFERENCES	20

I. INTRODUCTION

Transport coefficients, such as the self-diffusion coefficient D^* and ionic conductivity σ , are critical parameters for characterizing the performance of technologies such as batteries, fuel cells, and memristors. These coefficients are typically temperature dependent, with this temperature dependence commonly described by fitting empirical models to the measured data. Fitted models offer several benefits over raw data alone: they allow interpolation to reveal trends that otherwise would be obscured by noise, extrapolation of transport coefficients to outside the measured regime, and extraction of parameters, such as activation

*Contact author: andrew.mccluskey@bristol.ac.uk

†Contact author: b.j.morgan@bath.ac.uk

Published by the American Physical Society under the terms of the [Creative Commons Attribution 4.0 International](https://creativecommons.org/licenses/by/4.0/) license. Further distribution of this work must maintain attribution to the author(s) and the published article's title, journal citation, and DOI.

energies, that can be used to compare different materials or to explain microscopic dynamics [1,2].

The most widely used empirical model for describing the temperature dependence of transport coefficients is the Arrhenius equation. For diffusion coefficients, this takes the form

$$D^* = A \exp\left(\frac{-E_a}{RT}\right), \quad (1)$$

where E_a is the activation energy and A is the pre-exponential factor. For conductivity, the equivalent expression is [3]

$$\sigma T = A \exp\left(\frac{-E_a}{RT}\right). \quad (2)$$

Both equations can be linearized by taking logarithms, and it is common to plot the natural logarithm of D^* or σT against inverse temperature: this representation is known as an “Arrhenius plot,” and data that approximate a straight line are said to exhibit Arrhenius behavior.

While many systems show approximate Arrhenius behavior, this is not universal. For some systems, Arrhenius plots show pronounced curvature, corresponding to “non-Arrhenius” behavior that is not well modelled by the Arrhenius equation. In these cases, alternative models—such as the Vogel-Tammann-Fulcher equation [2,4–7]—may provide a better description. However, it is not always clear whether observed deviations indicate genuine non-Arrhenius behavior or simply reflect noise in the data, raising the question of which model is best supported by the data.

Fitting an empirical model requires determining the values of its parameters, e.g., for the Arrhenius model, the activation energy E_a and pre-exponential factor A . Because the underlying data have scatter or associated uncertainties, we need not only best-fit parameter values but also uncertainty estimates—known as inverse uncertainties—for those fitted parameters.

Once a model has been fitted, a common use is to extrapolate to temperatures outside the measured range; for example, to predict room-temperature conductivity from high-temperature simulation data [8,9]. Such extrapolations assume the model remains valid in the new regime. Additionally, for the extrapolated estimates to be meaningful, the uncertainties in the fitted parameters must propagate into the predicted values; without them, we have no way to judge how reliable the extrapolation is.

These challenges—obtaining fitted parameters with meaningful uncertainties, assessing which model is best supported by the data, and quantifying uncertainty in extrapolations—are nontrivial and often not addressed by standard fitting approaches. Bayesian methods provide a coherent framework for addressing all three: they yield full

probability distributions for fitted parameters, enable rational comparison of competing models, and allow uncertainty propagation into extrapolated predictions.

This tutorial discusses the application of Bayesian methods to the analysis of temperature-dependent transport data. Our examples focus on conductivity estimates from molecular dynamics simulations of superionic materials, but the methodology is generally applicable to any temperature-dependent data with associated uncertainties. The methods outlined here are available in the open-source Python package KINISI [10].

II. CHALLENGES IN TEMPERATURE-DEPENDENT MODELLING

Fitting empirical models to temperature-dependent transport data raises three significant challenges: selecting an appropriate model, determining parameter values with meaningful uncertainties, and extrapolating reliably to unmeasured temperatures. We examine each of these challenges below, before showing how Bayesian methods address them in Secs. III–V.

A. Model selection

Given a set of temperature-dependent data, how do we determine which empirical model best describes the underlying behavior? If we fit multiple candidate models and compare residuals, a more complex model with more free parameters will generally fit better than a simpler one, even when the additional complexity is not justified by the data. Using an overly complex model risks fitting noise rather than signal, producing parameter estimates that are precise but inaccurate, and extrapolations that are unreliable.

A model with more free parameters can fit a wider range of possible behaviors. If such a model fits the data, this is less informative than if a simpler, more constrained model fits equally well. The simpler model has made a more specific prediction that happened to be confirmed. When two models fit the data comparably, preferring the simpler model is not mere parsimony; it reflects the fact that the simpler model has been more strongly tested. What we need is a framework that balances fit quality against model complexity [11]. Bayesian model selection provides us with a rational approach to comparing competing models, and we discuss this in Sec. IV.

B. Parameter estimation

To fit a model to data, we must determine its parameter values. But given noisy or uncertain data, no single set of parameters is uniquely consistent with the observations—many different parameter combinations could plausibly describe the observed data. Rather than seeking a single “best fit,” we want a way to characterize the full set of

parameter values that are compatible with our data. Standard fitting approaches return point estimates, perhaps with symmetric error bars that quantify spread around the best fit. But this representation is incomplete; it does not capture the full range of models consistent with the data, nor any correlations or asymmetries in how parameters relate to one another. Bayesian posterior sampling addresses this by generating samples from the distribution of parameter values compatible with the observed data. This posterior distribution fully characterizes our uncertainty about the model parameters, as we discuss in Sec. III.

C. Extrapolation

Once we have fitted a temperature-dependent model, a common goal is to predict behavior at temperatures outside the measured range; for example, estimating room-temperature conductivity from high-temperature simulations [9]. For such extrapolations to be meaningful, the uncertainties in the fitted parameters must propagate into the extrapolated predictions. Without these, extrapolated values cannot be meaningfully compared to other values or assessed for reliability. Bayesian posterior sampling makes uncertainty propagation straightforward: the same parameter samples used to characterize the fit can be propagated through the model to generate a distribution of predicted values at any temperature, as we discuss in Sec. V.

III. PARAMETER ESTIMATION

Although model selection conceptually precedes fitting—we need to choose a model before fitting it—in practice, we must understand how to fit a single model before we can compare models. We therefore begin with parameter estimation. For illustration, we assume an Arrhenius model, and demonstrate how Bayesian posterior sampling yields parameter estimates with meaningful uncertainties. The approach is general, however, and applies to any parametric model.

Here, our goal is not to find a single “best” set of parameter values, but to characterize the full range of parameter values that are compatible with our data. For the Arrhenius model, this means finding all combinations of activation energy, E_a , and pre-exponential factor, A , that could plausibly describe the observed temperature-dependent conductivity.

Given a set of measured conductivities at different temperatures, how do we determine which parameter values are compatible with these data? The starting point is to ask, for a given choice of E_a and A , how well does the model predict the observed data? We quantify this using the likelihood: the probability of observing our data given a particular set of parameter values. This is a conditional probability, written $p(\mathbf{x} | \boldsymbol{\theta})$, where $\mathbf{x}$ denotes our data and $\boldsymbol{\theta}$ the parameters. If a parameter combination

produces model predictions that closely match the observations (within their uncertainties), the likelihood is high; if the predictions are far from the data, the likelihood is low.

The likelihood tells us how probable our data are given some parameters. But this is not quite what we want. We want the inverse: how probable are the parameters given our data? This is the posterior distribution, written $p(\boldsymbol{\theta} | \mathbf{x})$. Likelihood and posterior are related but not the same: $p(\mathbf{x} | \boldsymbol{\theta})$ is not the same as $p(\boldsymbol{\theta} | \mathbf{x})$. The relationship between them is given by Bayes’ theorem:

$$p(\boldsymbol{\theta} | \mathbf{x}) = \frac{p(\mathbf{x} | \boldsymbol{\theta}) \times p(\boldsymbol{\theta})}{p(\mathbf{x})}. \quad (3)$$

Here, $p(\boldsymbol{\theta} | \mathbf{x})$ is the posterior: the probability of parameters $\boldsymbol{\theta}$ given the data. $p(\mathbf{x} | \boldsymbol{\theta})$ is the likelihood. $p(\boldsymbol{\theta})$ is the prior: our knowledge about plausible parameter values before seeing the data. $p(\mathbf{x})$ is a normalizing constant that ensures the posterior integrates to one. Because $p(\mathbf{x})$ is a constant for any given dataset, the posterior is proportional to the likelihood times the prior:

$$p(\boldsymbol{\theta} | \mathbf{x}) \propto p(\mathbf{x} | \boldsymbol{\theta}) \times p(\boldsymbol{\theta}). \quad (4)$$

This means we can characterize the posterior by evaluating the likelihood and prior across parameter space.

The prior distribution, $p(\boldsymbol{\theta})$, encodes what we know about plausible parameter values before examining the data. In many cases, we have little prior knowledge, in which case we use broad, minimally informative priors: for example, a uniform distribution over a physically reasonable range. For the Arrhenius model, we require $E_a > 0$ (activation energies are positive) and $A > 0$ (the pre-exponential factor is positive). Beyond these physical constraints, we might specify uniform priors over ranges that comfortably encompass expected values. Provided the data are informative, the posterior will be dominated by the likelihood, so the precise choice of prior is often not critical; but it should always be reported.

For simple models with few parameters, we could evaluate the posterior on a fine grid across parameter space. But this becomes impractical as the number of parameters grows. Instead, we use sampling algorithms that explore parameter space efficiently, spending more time in regions of high posterior probability. Markov chain Monte Carlo (MCMC) is a widely used family of sampling algorithms. The details of how MCMC works are beyond the scope of this tutorial, but the essential idea is that it generates a sequence of parameter samples that, collectively, represent the posterior distribution. From these samples, we can compute summary statistics (means, medians, credible intervals) or visualize the full distribution.

To apply this framework, we first choose both a predictive model—for example, the Arrhenius equation—and a

statistical model for the uncertainties in our data. Together, these determine the likelihood function we use in Bayes' theorem. For a model with parameters θ , each choice of parameters predicts a conductivity value at each measured temperature; call these predictions $\mu(\theta)$. We also have our observed conductivity values, $\mathbf{x}$, which will not exactly match the predictions because of uncertainties in the data. A standard approach is to assume that the observed values are normally distributed around the model predictions. This assumption is often reasonable: if errors arise from many independent sources, their combined effect will tend toward a normal distribution. More fundamentally, if all we know about our errors is their mean and variance, the normal distribution is the only distribution that requires no further assumptions [12].

Under this assumption, our statistical model is a multivariate normal distribution, characterized by a covariance matrix Σ that describes the data uncertainties. The corresponding likelihood function is

$$p(\mathbf{x} | \theta) = (2\pi)^{-\frac{k}{2}} |\Sigma|^{-\frac{1}{2}} \exp \left\{ -\frac{1}{2} [\mu(\theta) - \mathbf{x}]^T \Sigma^{-1} [\mu(\theta) - \mathbf{x}] \right\}, \quad (5)$$

where k is the number of observations, $\mathbf{x}$ is the vector of observed conductivity values, $\mu(\theta)$ is the vector of model predictions, and Σ is the covariance matrix.

In practice, the procedure is as follows. We start with a model—here, the Arrhenius equation—and define priors for each parameter. We then construct the likelihood function using our measured data and their uncertainties. The sampler explores parameter space, evaluating the product of likelihood and prior at each point. Combinations that fit the data well and satisfy the prior receive high probability; those that fit poorly or lie outside the prior bounds receive low probability. After sufficient exploration, the collected samples represent the posterior distribution. From these samples, we can extract summary statistics (e.g., means, medians, credible intervals) or visualize the full distribution to understand correlations between parameters.

To illustrate this procedure, we use conductivity data from molecular dynamics simulations of a lithium-ion conductor, c-LLZO (see Methods for simulation details). The data comprise conductivity estimates at four temperatures, each with an associated uncertainty (see Methods for analysis details). To sample the posterior, we need to specify priors for E_a and A . Since we have no strong prior knowledge, we use broad uniform distributions that comfortably span the range of physically plausible values. With the likelihood and priors defined, we can run the sampler. Full details of the prior ranges and sampling procedure are given in the Methods.

The sampler returns a set of samples drawn from the joint posterior distribution $p(E_a, A | \mathbf{x})$. Each sample is a

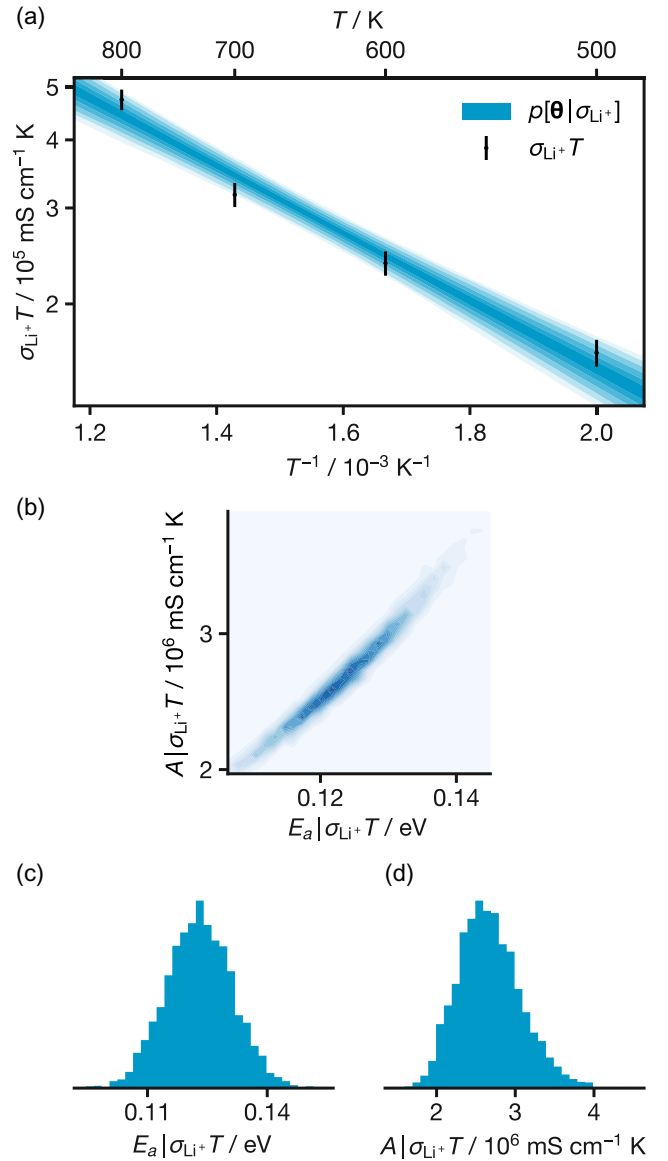


FIG. 1. (a) Lithium-ion conductivity in c-LLZO (black points; error bars show 95% credible intervals). Blue curves show Arrhenius models evaluated using parameter samples drawn from the posterior; darker shading indicates where more curves overlap. (b) Joint posterior distribution showing the correlation between E_a and A . (c),(d) Marginal posterior distributions for the activation energy and pre-exponential factor.

point in parameter space: a specific (E_a, A) combination that corresponds to a curve on the Arrhenius plot. Figure 1(a) shows the data with many such curves overlaid: the shaded bands indicate where curves cluster more or less densely, reflecting which parameter combinations are more probable given the data. Panel (b) shows these samples as a two-dimensional distribution, revealing that the parameters are highly correlated: higher values of E_a tend to be associated with higher values of A . Panels (c) and (d) show the marginal distributions for each parameter, obtained by

integrating the samples along each axis. This is equivalent to integrating the joint distribution over the other parameter, asking, for example, “what is the probability of this value of E_a , regardless of A ?”

The activation energy is approximately normally distributed, with mean and 95% credible interval of (0.123 ± 0.016) eV. The pre-exponential factor, however, is not normally distributed—it follows an approximately log-normal distribution. For this parameter, we report the median and 95% credible interval: $(2.626^{+0.904}_{-0.672}) \times 10^6 \text{ cm}^2 \text{ s}^{-1}$. This illustrates one advantage of Bayesian posterior sampling: we obtain the full distribution of each parameter, rather than assuming normality. For A , reporting a symmetric interval around the mean would misrepresent our uncertainty.

Arrhenius models are often fitted by taking logarithms and fitting a straight line to $\ln(\sigma T)$ versus $1/T$. This approach is convenient, but introduces bias: if the errors in σT are normally distributed, transforming to log space makes them log-normally distributed, violating the assumptions of linear regression and yielding biased estimates for E_a and A [13]. Our choice of a multivariate normal likelihood above also assumes normally distributed errors, so we must work with the untransformed data to satisfy this assumption.

IV. MODEL SELECTION

In Sec. III, we assumed the data followed the Arrhenius model. But how do we know this is the right choice? If we simply compare how well different models fit the data, a more complex model will generally fit better [14], even when the additional complexity is not justified. We need a rational approach that balances fit quality against model complexity. Bayesian model selection provides such a framework. Rather than asking “which model fits best?,” it asks “which model is most probable given the data?”; and in answering this question, it naturally penalizes unnecessary complexity.

To illustrate model selection, we compare the Arrhenius equation with the Vogel-Tammann-Fulcher (VTF) equation [4–6], an alternative model that allows the activation energy to vary with temperature. The VTF equation has the form:

$$\sigma T = A \exp \left[\frac{-B}{R(T - T_0)} \right], \quad (6)$$

where A is the pre-exponential factor, B is related to the activation energy, and T_0 is the Vogel temperature. The Arrhenius equation is a special case of VTF with $T_0 = 0$, so the VTF model will always fit the data at least as well as Arrhenius, because it has an additional parameter to adjust. The question is whether that additional parameter is justified by the data.

The Bayesian approach is to ask which model is more probable given the data; that is, compute $p(\text{model} | \mathbf{x})$, or in more compact notation, $p(m | \mathbf{x})$. Applying Bayes’ theorem at the level of models rather than parameters:

$$p(m | \mathbf{x}) \propto p(\mathbf{x} | m) \times p(m). \quad (7)$$

Here, $p(m)$ is our prior belief about the model before seeing the data. If we have no reason to prefer one model over another, we assign them equal prior probabilities. In that case, comparing posterior model probabilities reduces to comparing the marginal likelihood, $p(\mathbf{x} | m)$. The marginal likelihood is the probability of observing our data, averaged over all possible parameter values for that model. If we denote the parameters of model m as θ_m , this is,

$$p(\mathbf{x} | m) = \int \cdots \int_{\theta_m} p(\mathbf{x} | \theta_m, m) p(\theta_m | m) d\theta_m. \quad (8)$$

The first term in the integral is the likelihood: how well parameters θ_m predict the data. The second is the prior: how plausible those parameters are before seeing the data. This integral naturally penalizes complex models. A model with more parameters spreads its prior probability over a larger parameter space. Unless the data actually constrain those extra parameters, much of this space will have low likelihood, dragging down the integral. A simpler model spreads its prior over a smaller space; if the data fall within this space, the integral is correspondingly larger. This is the mathematical realization of the idea from Sec. II: a simpler model makes a more specific prediction. If that prediction is confirmed by the data, it provides stronger evidence than a flexible model that could have accommodated many different outcomes.

To compare two models, we compute the ratio of their marginal likelihoods:

$$B_{\beta\alpha} = \frac{p(\mathbf{x} | m_\beta)}{p(\mathbf{x} | m_\alpha)}. \quad (9)$$

This ratio is called the Bayes factor. If $B_{\beta\alpha} > 1$, the data favor model β ; if $B_{\beta\alpha} < 1$, they favor model α . The magnitude indicates the strength of the evidence. A common heuristic, due to Kass and Raftery, is that $\ln(B_{\beta\alpha}) > 5$ (corresponding to $B_{\beta\alpha} > 150$) constitutes “very strong” evidence for the preferred model [15].

The integral for the marginal likelihood can be challenging to compute, particularly for models with many parameters. Markov chain Monte Carlo, which we used in Sec. III for parameter estimation, does not directly provide the marginal likelihood. Nested sampling is an alternative algorithm designed specifically for this purpose: it efficiently estimates the marginal likelihood while also providing posterior samples as a byproduct. We do not detail the algorithm here, but point the reader to Sivia and Skilling for a pedagogical description [16].

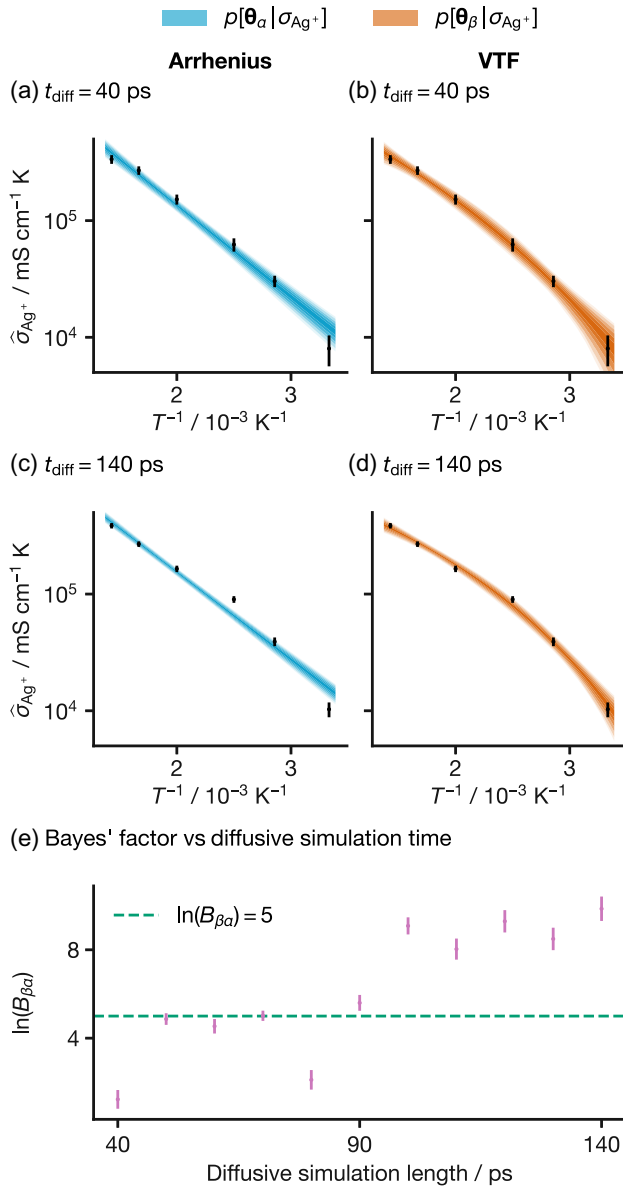


FIG. 2. Comparison of Arrhenius (a),(c) and VTF (b),(d) models for silver-ion conductivity in AgCrSe_2 . Black points show conductivity estimates with 95% credible intervals. Colored shading shows models evaluated at parameter values sampled from each posterior. Panels (a),(b) use data from short simulations (40-ps fitting windows; see Methods); panels (c),(d) use longer simulations (140 ps fitting windows). (e) Bayes factor comparing VTF to Arrhenius as a function of the fitting window width. The dashed line indicates $\ln(B_{\beta\alpha}) = 5$, the threshold for “very strong” evidence. Error bars represent a 95% confidence interval from 10 nested sampling repeats and reflect stochastic variability in the algorithm, not data uncertainty.

As an example of this approach, we consider silver-ion conductivity data from molecular dynamics simulations of AgCrSe_2 [2]. Figures 2(a) and 2(b) show ionic conductivities obtained from relatively short simulation trajectories (see Methods for details), fitted using the Arrhenius and

VTF models, respectively. Both models fit the data reasonably well, so how do we choose between them? The Bayes factor comparing the VTF model (β) to the Arrhenius model (α), assuming equal prior probabilities for each model, is $\ln(B_{\beta\alpha}) = 1.2 \pm 0.6$: this indicates weak evidence for the VTF model, but is well below the threshold for strong evidence. With this amount of data, we therefore have no compelling evidential basis for preferring the more complex VTF model.

Panels (c) and (d) show the same comparison using ionic conductivities from longer simulation trajectories. The conductivity data appear similar to those in panels (a) and (b), but the uncertainties are much smaller, and the deviation from linearity is more clearly resolved. The Bayes factor increases to $\ln(B_{\beta\alpha}) = 10.1 \pm 0.3$, well above the threshold of 5, providing strong evidence for the more complex VTF model over the simpler Arrhenius alternative.

Panel (e) shows how the Bayes factor evolves with increasing simulation trajectory lengths. As the conductivity data become more precise, the evidence for non-Arrhenius behavior grows. If the system exhibited true Arrhenius behavior, longer simulations would not increase the Bayes factor, and we would continue to prefer the simpler model.

This example illustrates two points. First, Bayesian model selection tells us whether a more complex model is supported by the data or if we should prefer a simpler alternative. Second, it tells us whether our data are sufficient to answer the question we are asking. With short simulation trajectories, we cannot confidently classify this behavior as non-Arrhenius; with longer trajectories, we can. More generally, this approach can be applied to any set of competing models; we need only compute the marginal likelihood for each.

V. EXTRAPOLATION

In Secs. III and IV, we fitted models to data within the measured temperature range. Often, however, we want to predict behavior at temperatures outside this range: for example, estimating room-temperature conductivity from high-temperature simulations. For such extrapolations to be meaningful, the uncertainties in the fitted parameters must propagate into the predicted values—a single extrapolated number without an associated uncertainty cannot be compared to other values or assessed for reliability.

Bayesian posterior sampling makes uncertainty propagation straightforward. In Sec. III, we obtained samples from the joint posterior distribution of parameters. Each sample represents a plausible (E_a, A) combination given the data. To extrapolate, we simply evaluate the model at the new temperature for each sample. The resulting collection of conductivity values represents the distribution

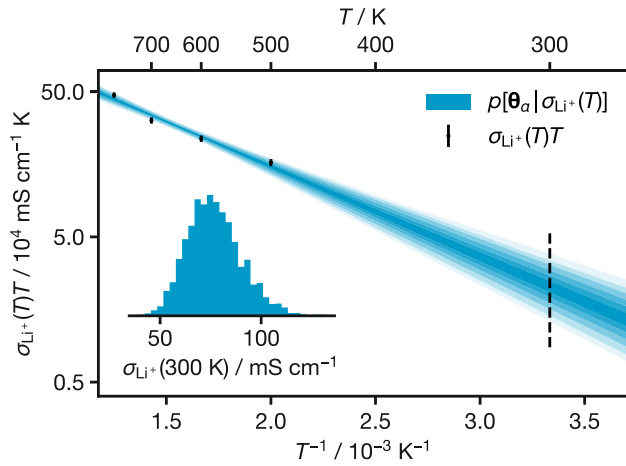


FIG. 3. Extrapolation of the LLZO conductivity model to 300 K (vertical dashed line). Blue shading shows Arrhenius models evaluated at parameter values sampled from the posterior (same data as Fig. 1). The shaded region widens at lower temperatures, reflecting increased uncertainty in the extrapolated predictions. Inset: distribution of predicted conductivity at 300 K.

of model predictions at that temperature; i.e., the range of conductivities compatible with our original data [17].

Returning to the LLZO example from Sec. III, we now extrapolate from the measured temperature range (500 K to 800 K) to 300 K, well outside the original data. For each posterior sample, we calculate $\sigma(300 \text{ K})$ using the Arrhenius equation. These predicted values form a distribution of conductivities at room temperature.

Figure 3 shows the result. The shaded bands in the main panel show the distribution of extrapolated conductivities at each temperature, and widen substantially as we move further from the measured range. The inset shows the distribution of predicted conductivity at 300 K. This distribution is broad and asymmetric, with a median and 95% credible interval of $(76.3^{+30.0}_{-23.0}) \text{ mS cm}^{-1}$. The breadth of this distribution illustrates why uncertainty propagation matters. A point estimate at 300 K would suggest a precise prediction; the full distribution reveals that the extrapolated conductivity spans nearly an order of magnitude. This uncertainty is not a failure of the method; it accurately reflects how much the data can tell us about behavior at temperatures we did not measure.

A caveat: for extrapolation to be meaningful, the fitted model must remain valid at the new temperature. Bayesian methods, however, cannot verify this assumption; if the transport mechanism changes—for example, if a different diffusion pathway becomes dominant at lower temperatures—the extrapolation will be misleading regardless of how carefully we propagate uncertainty.

VI. CONCLUSIONS

This tutorial has introduced Bayesian methods for analyzing temperature-dependent transport data, addressing

three common challenges: estimating model parameters with meaningful uncertainties, selecting between competing models, and extrapolating predictions to unmeasured temperatures. For parameter estimation, Bayesian posterior sampling returns the full distribution of parameter values compatible with the data, yielding meaningful uncertainties and capturing asymmetries and correlations that point estimates with symmetric error bars would miss. The approach requires only a model equation and appropriate priors, and applies to any parametric model. For model selection, Bayes' factors provide a rational basis for choosing between models of different complexity, based on the relative evidence provided by the data. For extrapolation, the same posterior samples enable straightforward uncertainty propagation, revealing how much (or how little) the data constrain predictions outside the measured range. The examples in this tutorial have focused on transport data from molecular dynamics simulations, but the same framework applies to any temperature-dependent measurements with associated uncertainties, such as reaction rates or mechanical properties. The methods described here are implemented in the open-source Python package KINISI [10], which we hope will make Bayesian analysis of transport data more broadly accessible.

VII. METHODS

All conductivity data used in this work were derived from molecular dynamics simulations. At each temperature, self-diffusion coefficients were estimated from the long-time slope of the mean-squared displacement (MSD) using approximate Bayesian regression [18], considering only the diffusive regime identified by visual inspection of log-log MSD plots. The range of lag times used for this regression is the fitting window. These were converted to conductivity estimates via the Nernst-Einstein relation, assuming a Haven ratio $H_R = 1$ [19]. All analysis was performed with KINISI-2.0.3 [10], using MDANALYSIS [20,21] for trajectory parsing; scripts are available in the Supplemental Material [22].

For $\text{Li}_7\text{La}_3\text{Zr}_2\text{O}_{12}$ (LLZO), classical molecular dynamics simulations were run using the METALWALLS code [23]. These simulations used the DIPPIM polarizable ion force field, as parameterized by Burbano *et al.* [24]. The cubic phase of LLZO was simulated in the NVT ensemble at temperatures of 500, 600, 700, and 800 K. Simulations were run for 50 ps using a 0.5-fs timestep and a Nosé-Hoover thermostat with a relaxation time of 121 fs [25–27]. The simulations used $2 \times 2 \times 2$ supercells with 1536 atoms, following the protocol in Ref. [24]. The simulation data are available under a CC BY 4.0 license on Zenodo [28]. Mean-squared displacements of the lithium ions were computed at lag times from 0.5 fs to 50 ps in intervals of 2.5 fs, with the start of the fitting window at 10 ps.

TABLE I. MSD fitting parameters for AgCrSe₂ at each simulation temperature. $\Delta t_{\min}$ is the minimum lag time used for fitting, corresponding to the onset of the diffusive regime. $\Delta(\Delta t)$ is the spacing between successive lag times in the MSD regression.

T/K	$\Delta t_{\min}/\text{ps}$	$\Delta(\Delta t)/\text{ps}$
300	70.0	3.00
350	50.0	1.25
400	30.0	0.60
500	25.0	0.30
600	22.5	0.25
700	20.0	0.20

The LLZO conductivity data were fitted to the Arrhenius equation using Markov chain Monte Carlo sampling. Parameter priors were uniform: $p(E_a) \sim \mathcal{U}(0, 0.5 \text{ eV})$ and $p(A) \sim \mathcal{U}(10^5, 10^7 \text{ cm s}^{-1})$. The posterior distribution was sampled using 32 walkers, generating 10 000 samples after a burn-in period of 500 samples. Samples were thinned by a factor of 10 to reduce autocorrelation. For extrapolation, posterior samples were used to estimate conductivity at $T = 300 \text{ K}$.

For AgCrSe₂, simulation trajectories were provided by the authors of Ref. [2]. These simulations covered temperatures of 300, 350, 400, 500, 600, and 700 K; full details are given in the original publication. The trajectories are available on Zenodo [29].

To investigate how simulation length affects the Bayes factor, we truncated the simulation trajectories to different lengths, setting $t_{\text{sim}} = \Delta t_{\min}(T) + \Delta t_{\text{diff}}$, where $\Delta t_{\min}(T)$ is the temperature-dependent onset of the diffusive regime (Table I). This gives a fitting window of width Δt_{diff} , which we varied from 40 to 140 ps [Fig. 2(e)]. At each temperature, MSDs were computed at lag times from $\Delta t_{\min}$ to t_{sim} —i.e., we always fit to the end of the available data—in increments of $\Delta(\Delta t)$ (Table I).

Bayesian model selection was performed using nested sampling [30] as implemented in the DYNesty Python package [31], with 500 live points and a stopping

criterion of 0.01. For the Arrhenius model, the parameter priors were $p(E_a) \sim \mathcal{U}(0, 1 \text{ eV})$ and $p(A) \sim \mathcal{U}(10^4, 10^8 \text{ cm}^2 \text{ s}^{-1})$. For the VTF model: $p(B) \sim \mathcal{U}(0, 1 \text{ eV})$, $p(A) \sim \mathcal{U}(10^4, 10^8 \text{ cm}^2 \text{ s}^{-1})$, and $p(T_0) \sim \mathcal{U}(0, 300 \text{ K})$.

ACKNOWLEDGMENTS

We thank the authors of the AgCrSe₂ work for providing simulation trajectories. This work used the Isambard 2 UK National Tier-2 HPC Service (<http://gw4.ac.uk/isambard/>) operated by GW4 and the UK Met Office and funded by EPSRC (Grant No. EP/T022078/1). The authors acknowledge the University of Bath's Research Computing Group for their support in running the LLZO molecular dynamics simulations. S.W.C. and B.J.M. acknowledge the support of the Faraday Institution through the CATMAT Project (Grant No. FIRG016). B.J.M. acknowledges support from the Royal Society (Grants No. UF130329 and No. URF\R\191006). This work has benefitted from the input of all members of the KINISI development team.

A.R.M.: conceptualization, formal analysis, investigation, methodology, software, visualization, writing—original draft. S.W.C.: methodology, resources, writing—review and editing. B.J.M.: conceptualization, methodology, software, writing—review and editing.

DATA AVAILABILITY

The data that support the findings of this article are openly available [22,28,29].

APPENDIX A: MEAN-SQUARED DISPLACEMENT OF LLZO AS A FUNCTION OF TEMPERATURE

Figure 4 presents the mean-squared displacement data and resulting posterior distribution of linear models for a range of temperatures for LLZO. The marginal posterior distribution for D^* is then used to find σ , which is included in the LLZO analysis in the main text (Figs. 1 and 3).

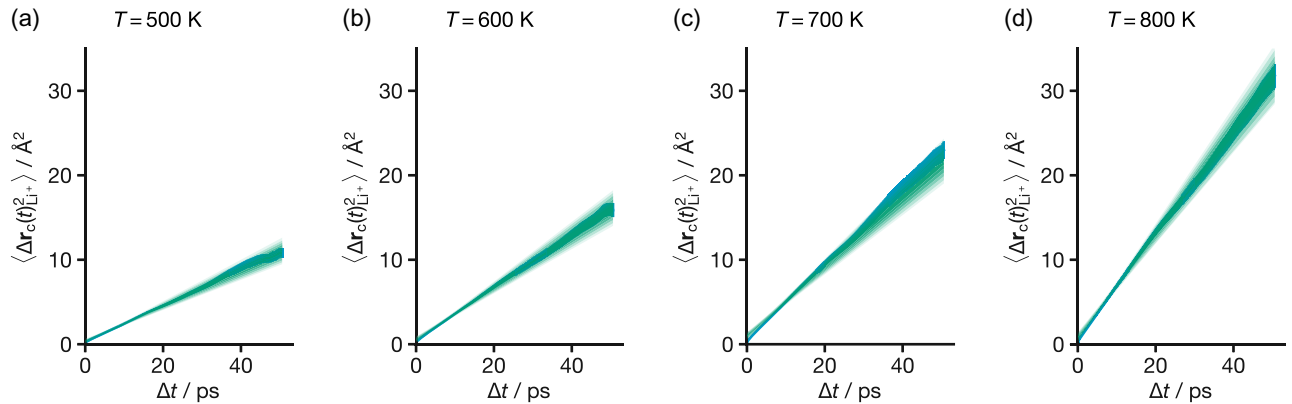


FIG. 4. The mean squared displacement data and associated distribution of linear models at temperatures of (a) 500 K, (b) 600 K, (c) 700 K, and (d) 800 K with 50 ps of LLZO simulation.

APPENDIX B: TEMPERATURE DEPENDENCE AS A FUNCTION OF DIFFUSIVE SIMULATION TIME FOR AgCrSe₂

Figures 5 to 15 show the mean-squared displacement as a function of lag time for temperatures of 300, 350,

400, 500, 600, and 700 K, for increasing total simulation lengths. Longer simulations provide more data in the diffusive regime, reducing the uncertainty in the conductivity estimates and, as discussed in Sec. IV, increasing the Bayesian evidence for the VTF model.

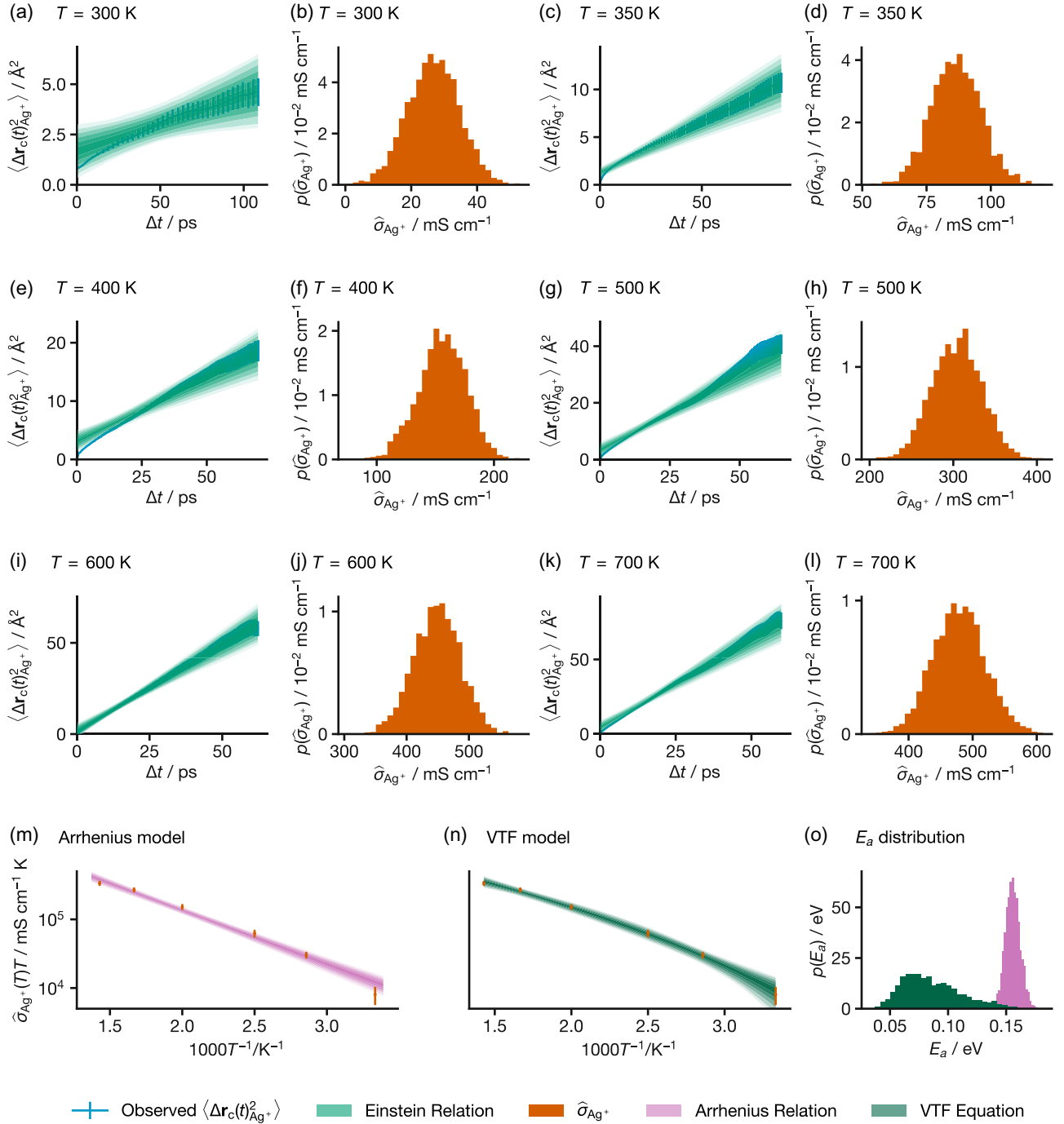


FIG. 5. The mean squared displacement data and associated σ distributions at temperatures of 300, 350, 400, 500, 600, and 700 K with 40 ps of diffusive simulation (a)–(l), the appropriate Arrhenius (m) and VTF model (n) plots and the resulting distributions of E_a from each modelling approach (o).

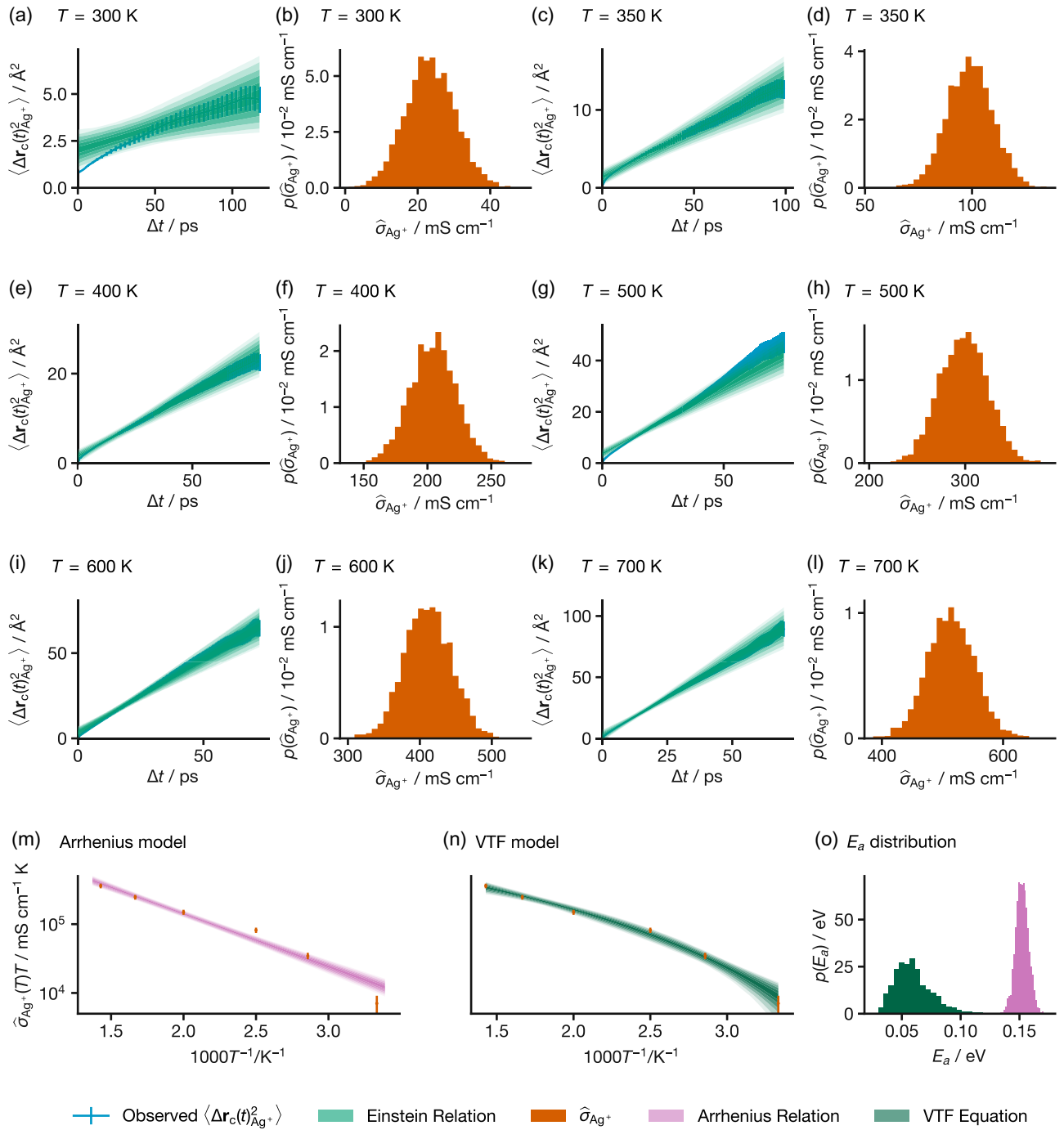


FIG. 6. The mean squared displacement data and associated σ distributions at temperatures of 300, 350, 400, 500, 600, and 700 K with 50 ps of diffusive simulation (a)–(l), the appropriate Arrhenius (m) and VTF model (n) plots and the resulting distributions of E_a from each modelling approach (o).

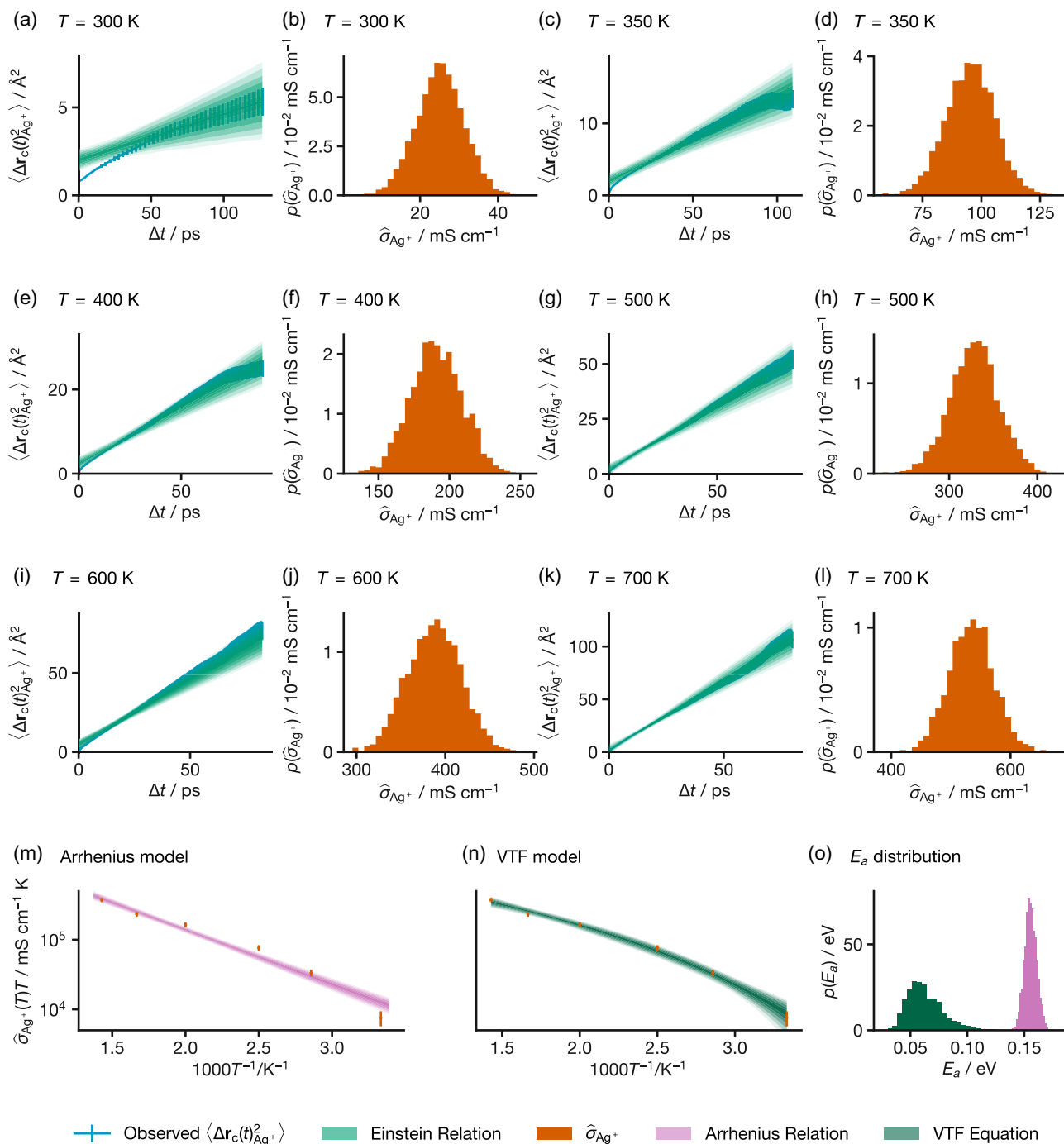


FIG. 7. The mean squared displacement data and associated σ distributions at temperatures of 300, 350, 400, 500, 600, and 700 K with 60 ps of diffusive simulation (a)–(l), the appropriate Arrhenius (m) and VTF model (n) plots and the resulting distributions of E_a from each modelling approach (o).

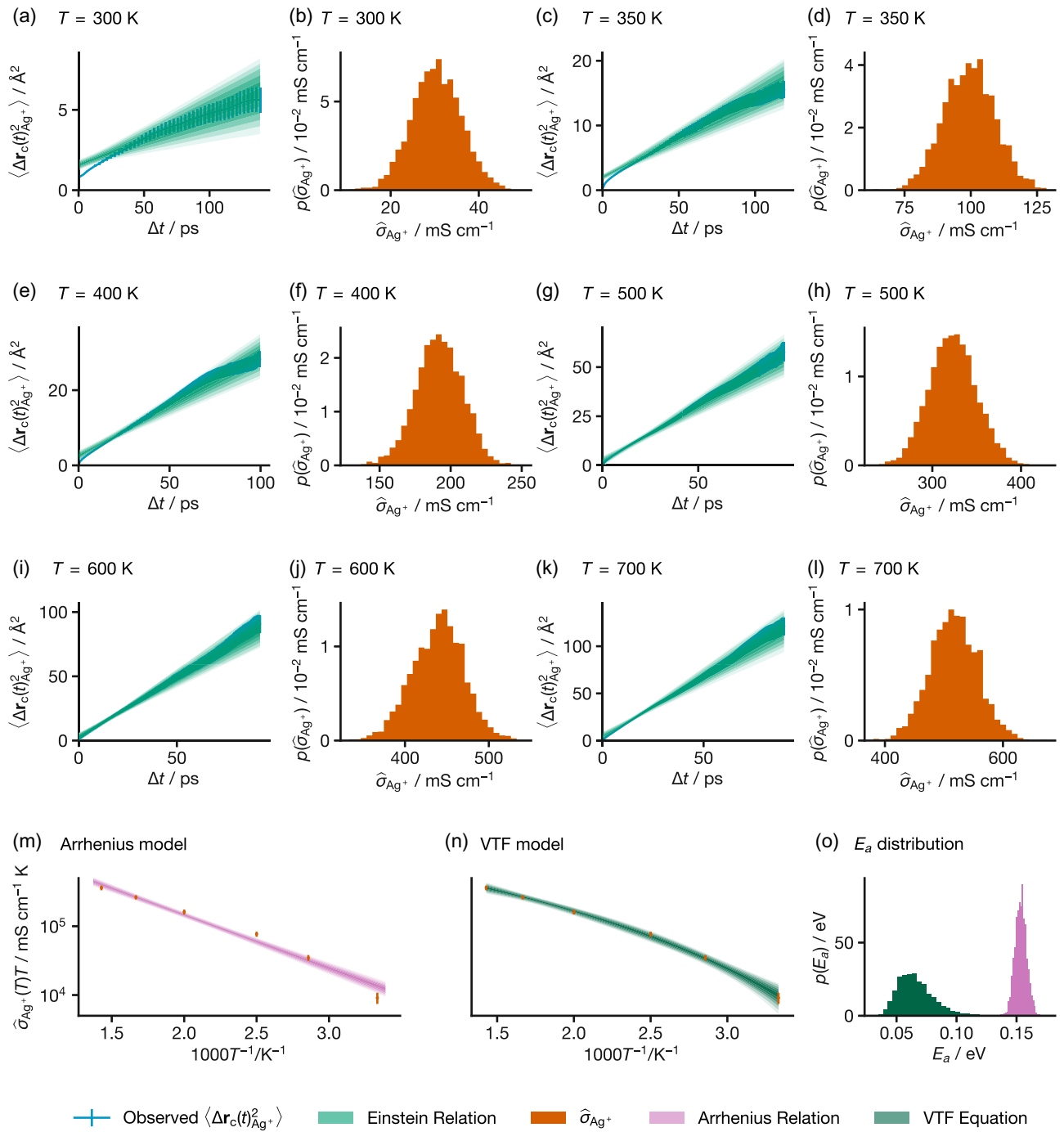


FIG. 8. The mean squared displacement data and associated σ distributions at temperatures of 300, 350, 400, 500, 600, and 700 K with 70 ps of diffusive simulation (a)–(l), the appropriate Arrhenius (m) and VTF model (n) plots and the resulting distributions of E_a from each modelling approach (o).

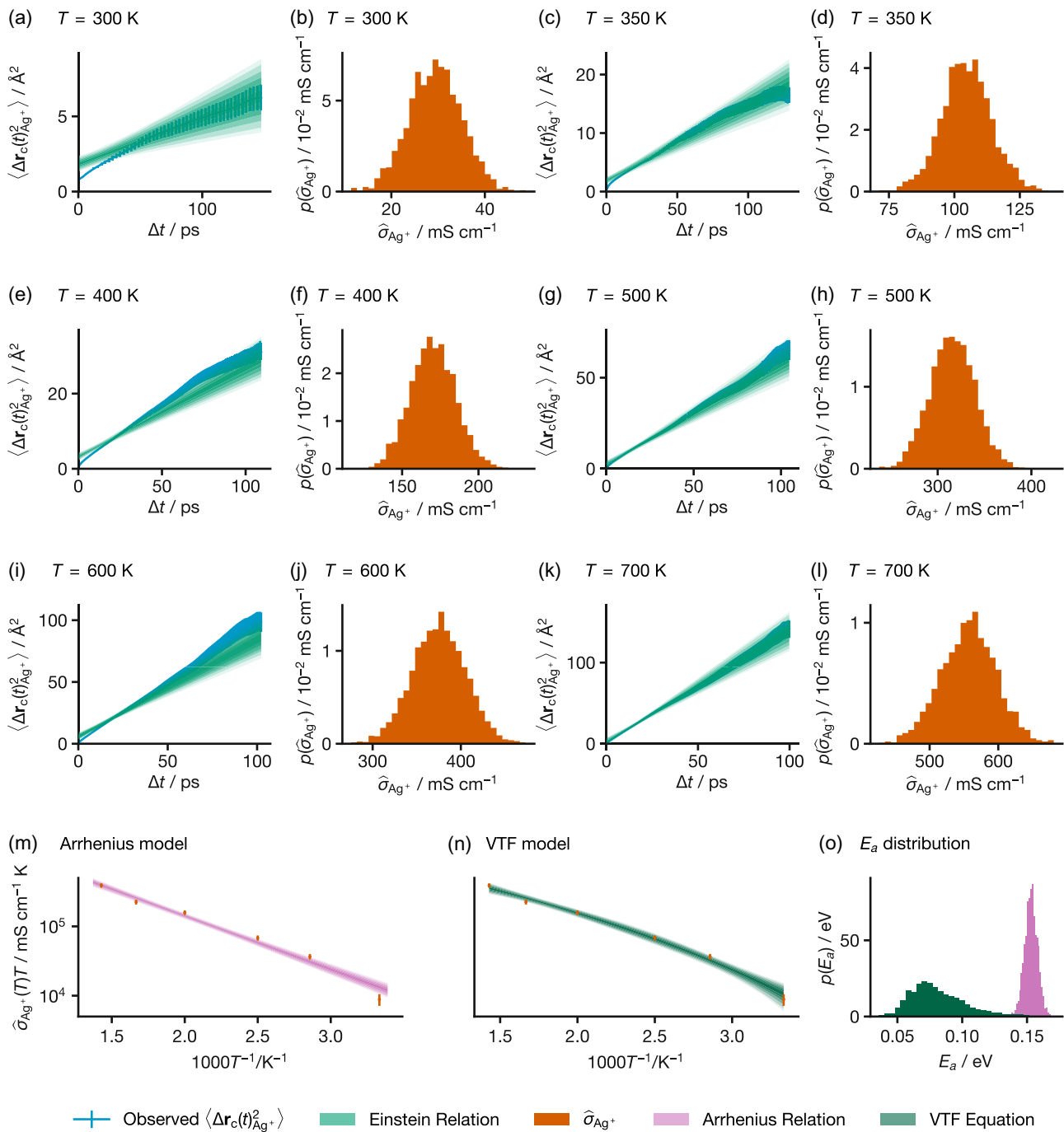


FIG. 9. The mean squared displacement data and associated σ distributions at temperatures of 300, 350, 400, 500, 600, and 700 K with 80 ps of diffusive simulation (a)–(l), the appropriate Arrhenius (m) and VTF model (n) plots and the resulting distributions of E_a from each modelling approach (o).

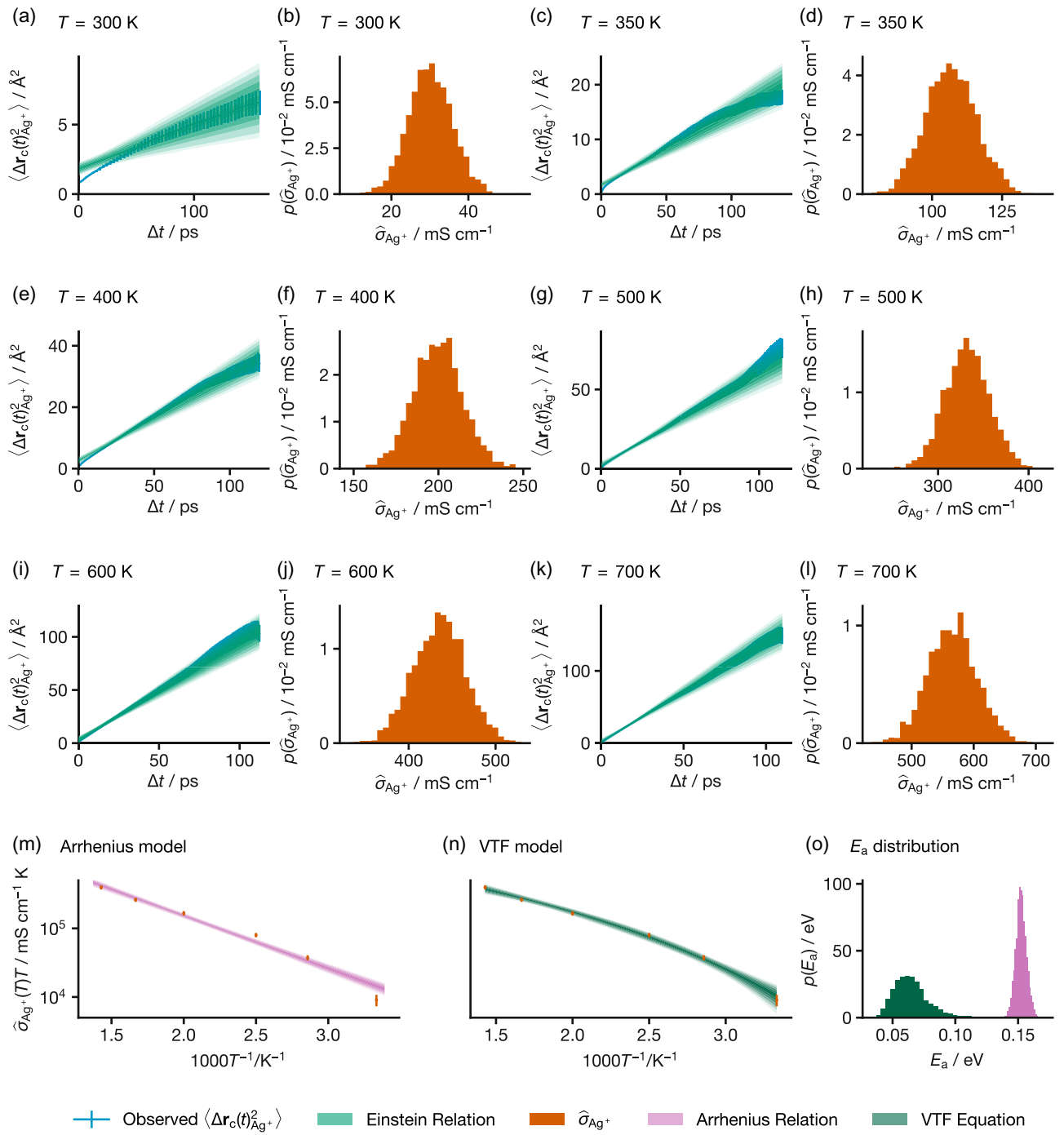


FIG. 10. The mean squared displacement data and associated σ distributions at temperatures of 300, 350, 400, 500, 600, and 700 K with 90 ps of diffusive simulation (a)–(l), the appropriate Arrhenius (m) and VTF model (n) plots and the resulting distributions of E_a from each modelling approach (o).

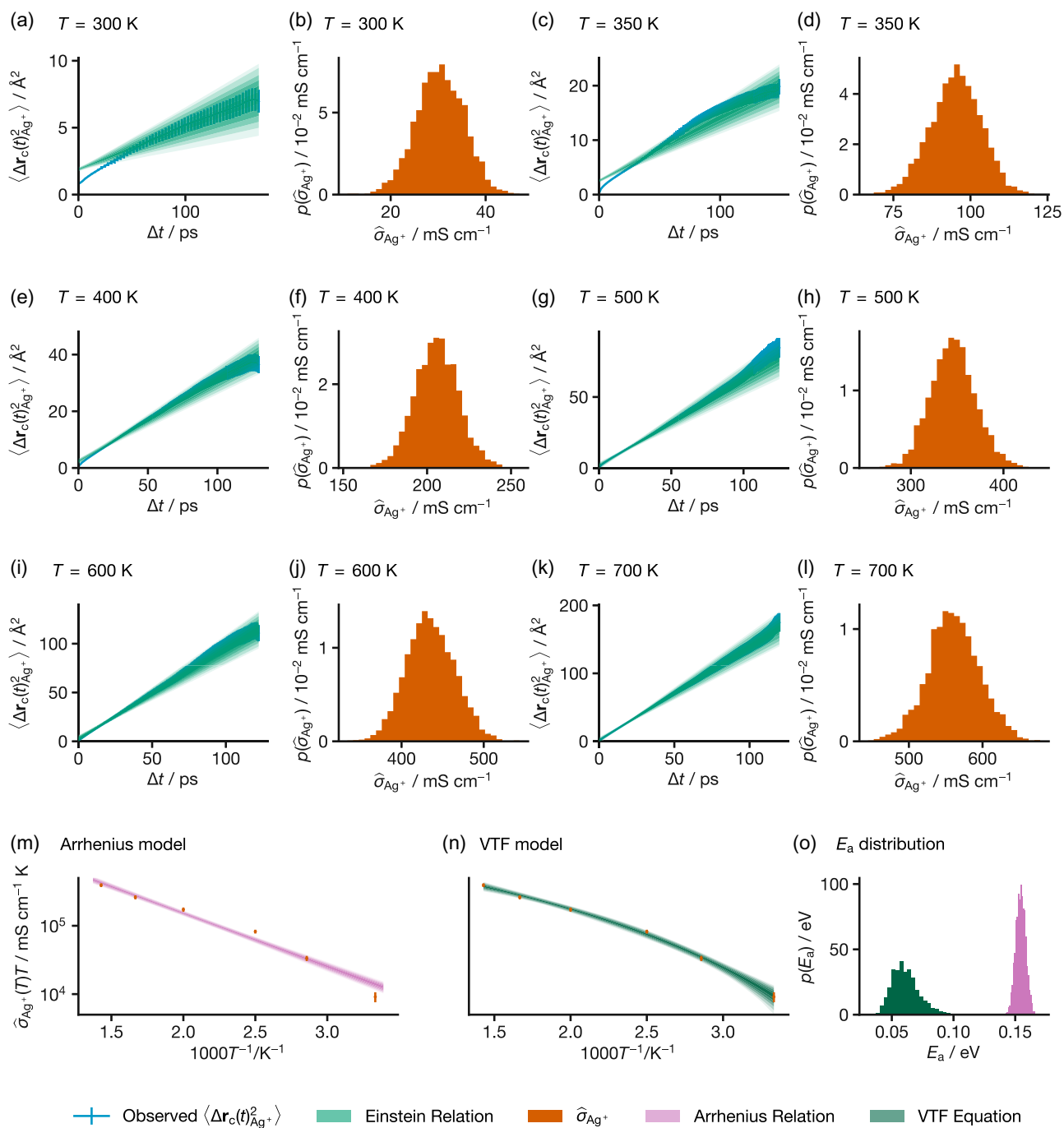


FIG. 11. The mean squared displacement data and associated σ distributions at temperatures of 300, 350, 400, 500, 600, and 700 K with 100 ps of diffusive simulation (a)–(l), the appropriate Arrhenius (m) and VTF model (n) plots and the resulting distributions of E_a from each modelling approach (o).

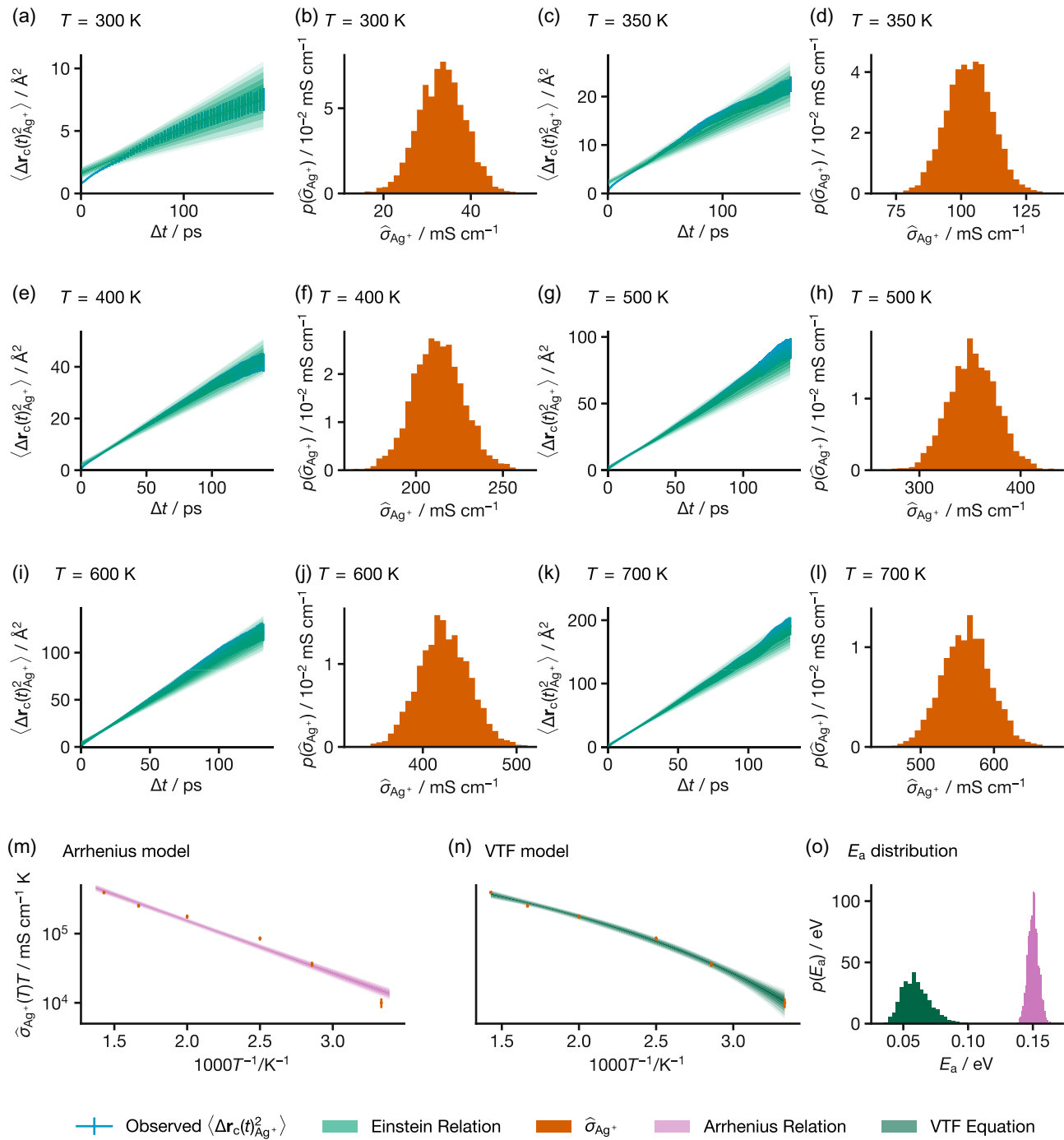


FIG. 12. The mean squared displacement data and associated σ distributions at temperatures of 300, 350, 400, 500, 600, and 700 K with 110 ps of diffusive simulation (a)–(l), the appropriate Arrhenius (m) and VTF model (n) plots and the resulting distributions of E_a from each modelling approach (o).

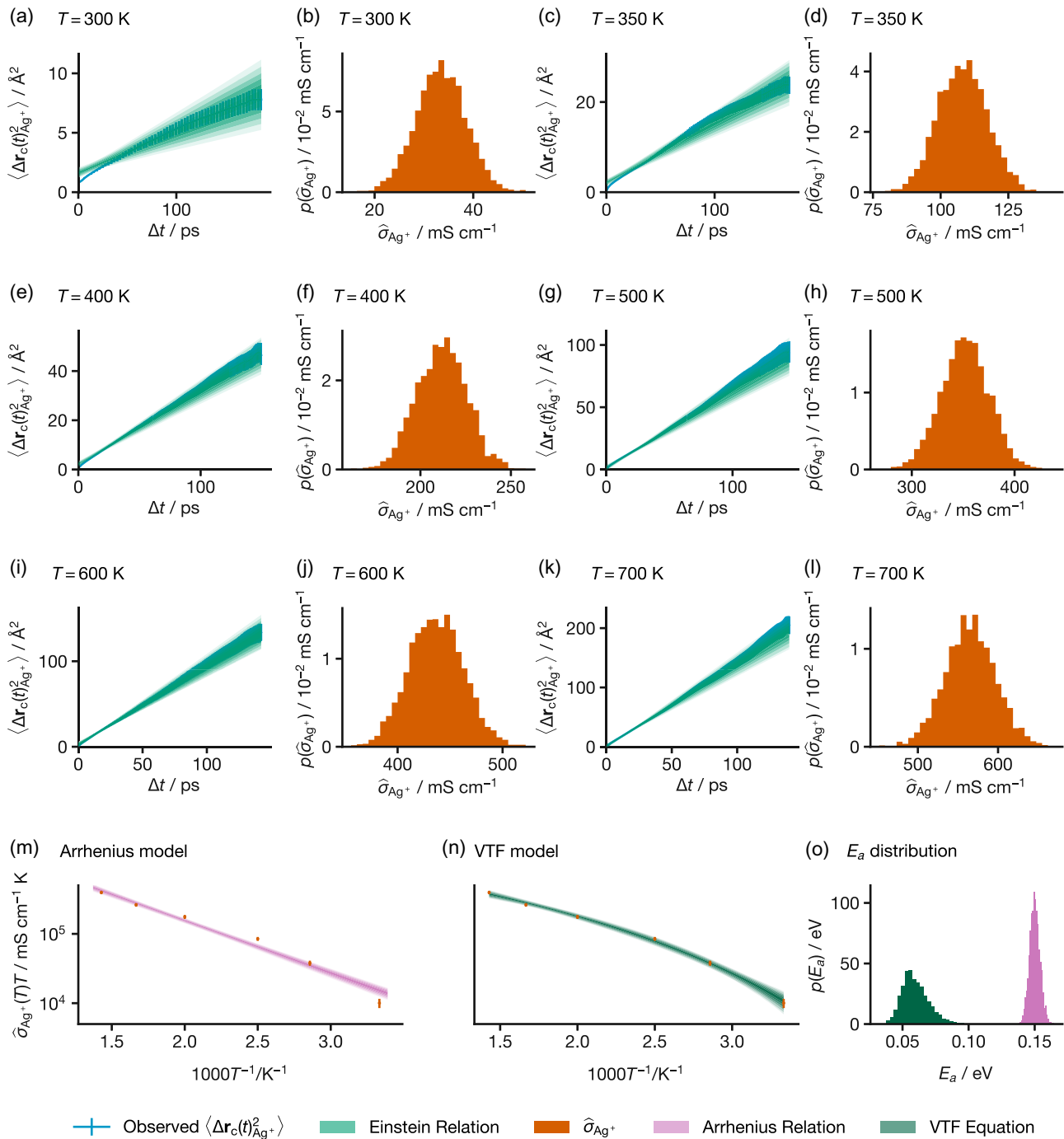


FIG. 13. The mean squared displacement data and associated σ distributions at temperatures of 300, 350, 400, 500, 600, and 700 K with 120 ps of diffusive simulation (a)–(l), the appropriate Arrhenius (m) and VTF model (n) plots and the resulting distributions of E_a from each modelling approach (o).

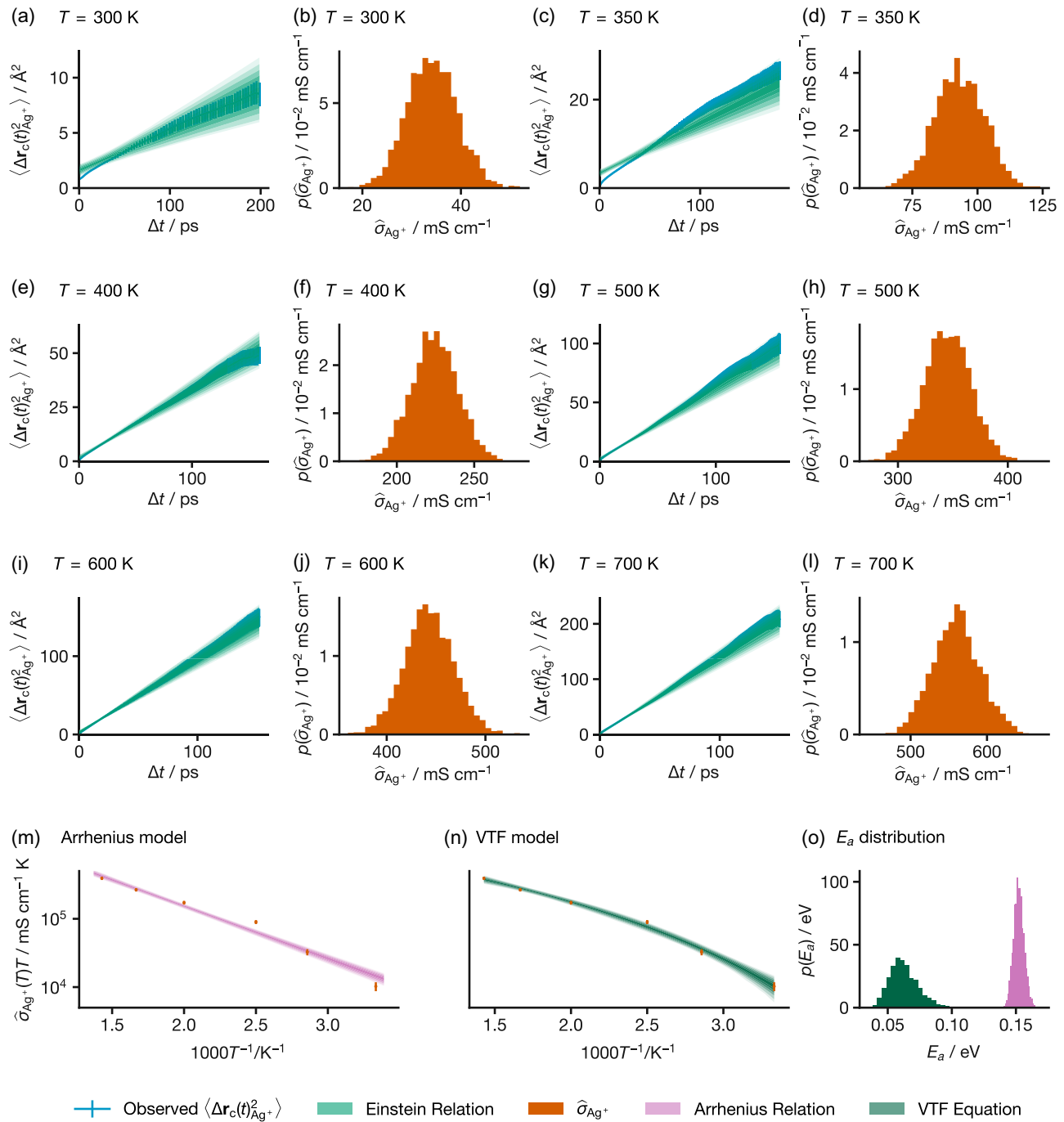


FIG. 14. The mean squared displacement data and associated σ distributions at temperatures of 300, 350, 400, 500, 600, and 700 K with 130 ps of diffusive simulation (a)–(l), the appropriate Arrhenius (m) and VTF model (n) plots and the resulting distributions of E_a from each modelling approach (o).

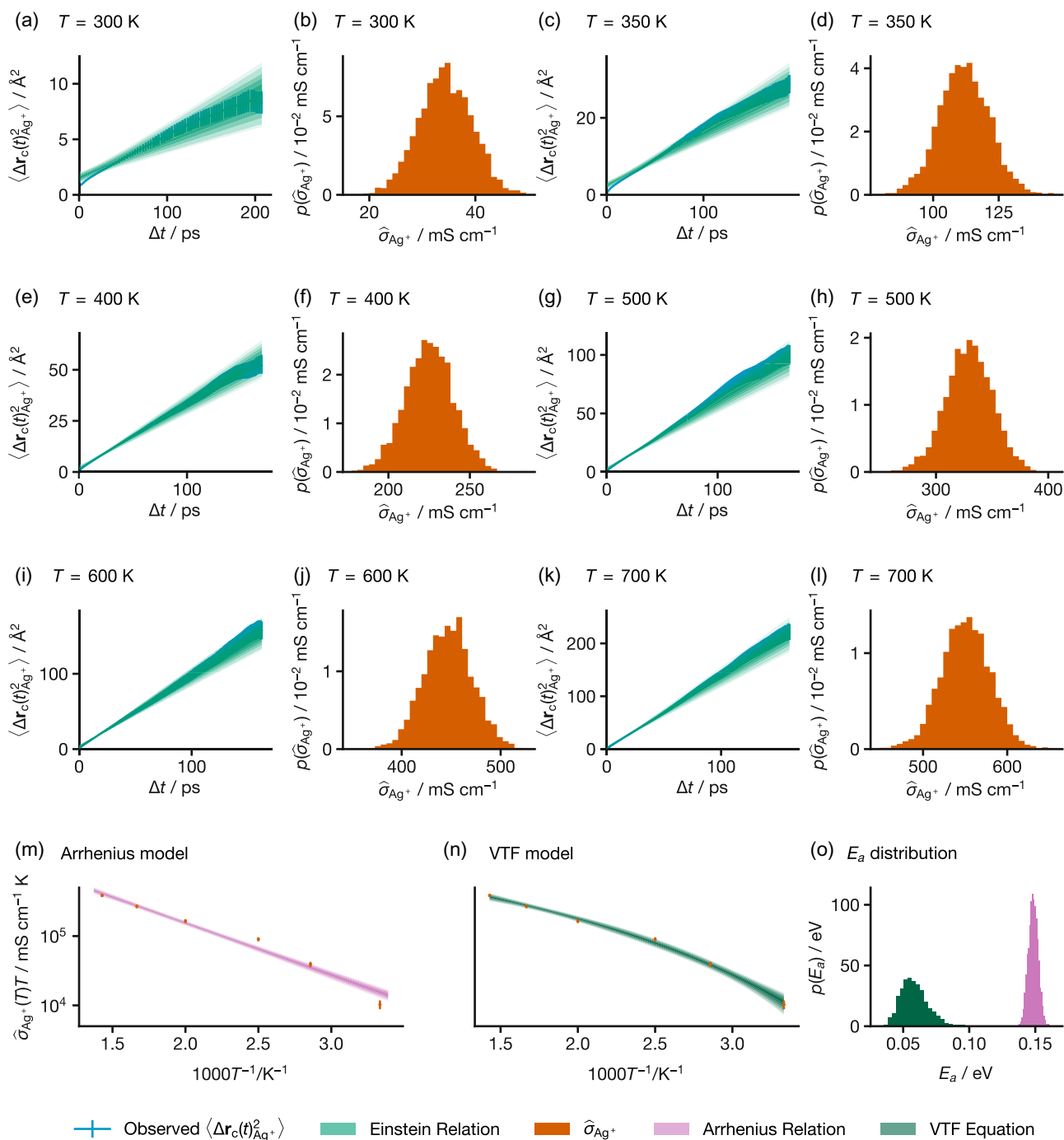


FIG. 15. The mean squared displacement data and associated σ distributions at temperatures of 300, 350, 400, 500, 600, and 700 K with 140 ps of diffusive simulation (a)–(l), the appropriate Arrhenius (m) and VTF model (n) plots and the resulting distributions of E_a from each modelling approach (o).

- [1] P. Canepa, Pushing forward simulation techniques of ion transport in ion conductors for energy materials, *ACS Mater. Au* **3**, 75 (2023).
- [2] J. Wang, J. Ding, O. Delaire, and G. Arya, Atomistic mechanisms underlying non-Arrhenius ion transport in superionic conductor AgCrSe₂, *ACS Appl. Energy Mater.* **4**, 7157 (2021).
- [3] The Nernst-Einstein relation between σ and D^* introduces a T^{-1} factor, so the Arrhenius form applies to σT rather than σ .
- [4] H. Vogel, Das Temperaturabhaengigkeitsgesetz der Viskositaet von Fluessigkeiten, *Phys. Z.* **22**, 645 (1921).
- [5] G. S. Fulcher, Analysis of recent measurements of the viscosity of glasses, *J. Am. Ceram. Soc.* **8**, 339 (1925).
- [6] G. Tammann and W. Hesse, Die Abhängigkeit der Viskosität von der Temperatur bei unterkühlten Flüssigkeiten, *Z. Anorg. Allg. Chem.* **156**, 245 (1926).
- [7] V. K. de Souza and D. J. Wales, Super-Arrhenius diffusion in an undercooled binary Lennard-Jones liquid results from a quantifiable correlation effect, *Phys. Rev. Lett.* **96**, 057802 (2006).
- [8] S. R. Yeandel, D. O. Scanlon, and P. Goddard, Enhanced Li-ion dynamics in trivalently doped lithium phosphidosilicate Li₂SiP₂: A candidate material as a solid Li electrolyte, *J. Mater. Chem. A* **7**, 3953 (2019).
- [9] B. A. Goldmann, C. Rosenbach, H. A. Evans, B. Helm, B. Wankmiller, O. Maus, E. Suard, L. F. Nazar, M. R. Hansen, B. J. Morgan, M. S. Islam, and W. G. Zeier, Rotational stacking faults in the ionic conductor Li₃ScCl₆, *Chem. Mater.* **37**, 9858 (2025).
- [10] A. R. McCluskey, A. G. Squires, J. Dunn, S. W. Coles, and B. J. Morgan, Kinisi: Bayesian analysis of mass transport from molecular dynamics simulations, *J. Open Source Softw.* **9**, 5984 (2024).
- [11] Information criteria such as AIC and BIC address this problem [32,33], but rely on assumptions such as normally distributed parameters that may not hold in practice.
- [12] This follows from the principle of maximum entropy [34].
- [13] A. R. McCluskey, Is there still a place for linearization in the chemistry curriculum? *J. Chem. Educ.* **100**, 4174 (2023).
- [14] J. Mayer, K. Khairy, and J. Howard, Drawing an elephant with four complex parameters, *Am. J. Phys.* **78**, 648 (2010).
- [15] R. E. Kass and A. E. Raftery, Bayes factors, *J. Am. Stat. Assoc.* **90**, 773 (1995).
- [16] D. S. Sivia and J. Skilling, *Data Analysis: A Bayesian Tutorial* (Oxford University Press, Oxford, UK, 2006), 2nd ed.
- [17] This extrapolated distribution captures parameter uncertainty but not estimated measurement noise. For extrapolation far from the measured range, parameter uncertainty typically dominates, so this distinction is rarely important in practice.
- [18] A. R. McCluskey, S. W. Coles, and B. J. Morgan, Accurate estimation of diffusion coefficients and their uncertainties from computer simulation, *J. Chem. Theory Comput.* **21**, 79 (2025).
- [19] A temperature-independent Haven ratio would uniformly scale the conductivity values without affecting the temperature dependence; a temperature-dependent ratio would introduce additional curvature, but the analysis procedure would be unchanged.
- [20] N. Michaud-Agrawal, E. J. Denning, T. B. Woolf, and O. Beckstein, MDAnalysis: A toolkit for the analysis of molecular dynamics simulations, *J. Comput. Chem.* **32**, 2319 (2011).
- [21] R. Gowers, M. Linke, J. Barnoud, T. Reddy, M. Melo, S. Seyler, J. Domański, D. Dotson, S. Buchoux, I. Kenney, and O. Beckstein, in *Proceedings of the 15th Python in Science Conference* (Austin, US, 2016), pp. 98–105.
- [22] A. R. McCluskey, S. W. Coles, and B. J. Morgan, model-arrhenius-1.0.0, 2025, <https://github.com/scams-research/model-arrhenius>.
- [23] A. Marin-Lafleche, M. Haeefe, L. Scalfi, A. Coretti, T. Dufils, G. Jeanmairet, S. Reed, A. Serva, R. Berthin, C. Bacon, S. Bonella, B. Rotenberg, P. Madden, and M. Salanne, MetalWalls: A classical molecular dynamics software dedicated to the simulation of electrochemical systems, *J. Open Source Softw.* **5**, 2373 (2020).
- [24] M. Burbano, D. Carlier, F. Boucher, B. J. Morgan, and M. Salanne, Sparse cyclic excitations explain the low ionic conductivity of stoichiometric Li₇La₃Zr₂O₁₂, *Phys. Rev. Lett.* **116**, 135901 (2016).
- [25] S. Nosé, A unified formulation of the constant temperature molecular dynamics methods, *J. Chem. Phys.* **81**, 511 (1984).
- [26] W. G. Hoover, Canonical dynamics: Equilibrium phase-space distributions, *Phys. Rev. A* **31**, 1695 (1985).
- [27] G. J. Martyna, M. L. Klein, and M. Tuckerman, Nosé-Hoover chains: The canonical ensemble via continuous dynamics, *J. Chem. Phys.* **97**, 2635 (1992).
- [28] Samuel W. Coles, LLZO Simulation Trajectories, 2025, <https://doi.org/10.5281/zenodo.17702491>.
- [29] J. Wang, J. Ding, O. Delaire, and G. Arya, AgCrSe₂ Simulation Trajectories, 2025, <https://doi.org/10.5281/zenodo.17910336>.
- [30] J. Skilling, in *AIP Conference Proceedings* (AIP, Garching, DE, 2004), Vol. 735, pp. 395–405.
- [31] J. S. Speagle, DYNESTY: A dynamic nested sampling package for estimating Bayesian posteriors and evidences, *Mon. Not. R. Astron. Soc.* **493**, 3132 (2020).
- [32] H. Akaike, A new look at the statistical model identification, *IEEE Trans. Autom. Control* **19**, 716 (1974).
- [33] G. Schwarz, Estimating the dimension of a model, *Ann. Stat.* **6**, 461 (1978).
- [34] E. T. Jaynes, Information theory and statistical mechanics, *Phys. Rev.* **106**, 620 (1957).