



Delft University of Technology

Green AI in Action

Strategic Model Selection for Ensembles in Production

Nijkamp, Nienke; Sallou, June; van der Heijden, Niels; Cruz, Luís

DOI

[10.1145/3664646.3664763](https://doi.org/10.1145/3664646.3664763)

Publication date

2024

Document Version

Final published version

Published in

Alware 2024

Citation (APA)

Nijkamp, N., Sallou, J., van der Heijden, N., & Cruz, L. (2024). Green AI in Action: Strategic Model Selection for Ensembles in Production. In *Alware 2024: Proceedings of the 1st ACM International Conference on AI-Powered Software* (pp. 50-58). ACM. <https://doi.org/10.1145/3664646.3664763>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.



Green AI in Action: Strategic Model Selection for Ensembles in Production

Nienke Nijkamp

Delft University of Technology, Deloitte Risk Advisory BV
The Netherlands
n.nijkamp@student.tudelft.nl

Niels van der Heijden

University of Amsterdam, Deloitte Risk Advisory BV
The Netherlands
n.vanderheijden@uva.nl

June Sallou

Delft University of Technology
The Netherlands
J.Sallou@tudelft.nl

Luís Cruz

Delft University of Technology
The Netherlands
L.Cruz@tudelft.nl

ABSTRACT

Integrating Artificial Intelligence (AI) into software systems has significantly enhanced their capabilities while escalating energy demands. Ensemble learning, combining predictions from multiple models to form a single prediction, intensifies this problem due to cumulative energy consumption.

This paper presents a novel approach to model selection that addresses the challenge of balancing the accuracy of AI models with their energy consumption in a live AI ensemble system. We explore how reducing the number of models or improving the efficiency of model usage within an ensemble during inference can reduce energy demands without substantially sacrificing accuracy.

This study introduces and evaluates two model selection strategies, *Static* and *Dynamic*, for optimizing ensemble learning systems' performance while minimizing energy usage. Our results demonstrate that the *Static* strategy improves the F1 score beyond the baseline, reducing average energy usage from 100% from the full ensemble to 62%. The *Dynamic* strategy further enhances F1 scores, using on average 76% compared to 100% of the full ensemble.

Moreover, we propose an approach that balances accuracy with resource consumption, significantly reducing energy usage without substantially impacting accuracy. This method decreased the average energy usage of the *Static* strategy from approximately 62% to 14%, and for the *Dynamic* strategy, from around 76% to 57%.

Our field study of Green AI using an operational AI system developed by a large professional services provider shows the practical applicability of adopting energy-conscious model selection strategies in live production environments.

CCS CONCEPTS

• **Software and its engineering** → **Software creation and management**; • **Computing methodologies** → **Artificial intelligence**; • **Applied computing** → Document management and text processing.



This work is licensed under a Creative Commons Attribution 4.0 International License.

AIware '24, July 15–16, 2024, Porto de Galinhas, Brazil

© 2024 Copyright held by the owner/author(s).

ACM ISBN 979-8-4007-0685-1/24/07

<https://doi.org/10.1145/3664646.3664763>

KEYWORDS

Green AI, Ensemble Learning, Model Selection, Information Extraction

ACM Reference Format:

Nienke Nijkamp, June Sallou, Niels van der Heijden, and Luís Cruz. 2024. Green AI in Action: Strategic Model Selection for Ensembles in Production. In *Proceedings of the 1st ACM International Conference on AI-Powered Software (AIware '24)*, July 15–16, 2024, Porto de Galinhas, Brazil. ACM, New York, NY, USA, 9 pages. <https://doi.org/10.1145/3664646.3664763>

1 INTRODUCTION

Over recent years, the integration of Artificial Intelligence (AI) into modern software systems, commonly known as AIware or AI-powered software [32], has experienced significant growth [19]. As a result, the demand for resources, particularly the energy needed for training, deploying, and running inferences with these AI models, has surged considerably [2, 8, 42]. To illustrate the resource consumption associated with inference, a ChatGPT-like application handling 11 million requests per hour is estimated to emit 12,800 tons of CO₂ annually, making inference 25 times more carbon-intensive than training GPT-3 [8].

Ensemble learning, a method that combines multiple models to create a more effective solution, while highly effective, tends to intensify this energy consumption problem due to the cumulative energy requirements of the individual models [11, 28].

The focus of the AI research community has predominantly been on improving the accuracy of AI models, overlooking the significant energy costs associated with them. AI's rising environmental and financial cost has led to a pressing need for a more balanced approach to AI development that considers accuracy and energy efficiency [42]. The emerging field of Green AI addresses this gap, promoting a favourable trade-off between efficiency and accuracy [39]. A significant gap remains in implementing Green AI principles within the industry [44]. To bridge this gap, we have partnered with Deloitte NL [34], a large professional services network specializing in, amongst others, digital risk solutions, to carry out a field experiment on an active AI system.

Our research explores the intersection of Green AI and ensemble learning in a production environment, a domain that has yet to be thoroughly investigated. Simply running all models every time is not an efficient strategy [49]. The challenge is finding a more intelligent approach for running inference with ensemble

learning [46], reducing energy consumption while maintaining or improving accuracy.

In this paper, we propose a solution that involves a selective approach to using models within an ensemble. The core of this approach is the concept of *model selection strategies*, which refers to various methods of selecting specific subsets of models for individual tasks rather than using the complete set of models for every task [6]. Our goal in implementing these model selection strategies is to balance achieving accuracy and managing computational costs. This goal leads to the following research question.

Research Question: What are the impacts of implementing model selection strategies on the accuracy and energy usage of ensemble learning systems?

Two model selection strategies, *Static* and *Dynamic*, are investigated for optimizing model performance. Both strategies start by evaluating all subsets of the entire ensemble on a validation set, selecting the combination with the highest accuracy. *Static* selection identifies the best overall model selection, while *Dynamic* selection chooses the best selection per specific property within the domain.

Additionally, we consider the computational cost per model by employing a metric that discounts accuracy with energy consumption. This method is incorporated into both the *Static* and *Dynamic* selection strategies, ensuring a balanced consideration of performance and resource efficiency.

This empirical study explores the potential of model selection strategies within an ensemble of AI models. We conduct a field study on DocQMiner [12], a live, industry-used AI system for text information extraction from large volumes of documents, to assess the practical implications of our proposed model selection strategies.

This research provides practical insights into making model ensembles more efficient by examining and implementing our model selection strategies within an ensemble learning context. Our contributions to the field of AI in software systems using ensemble learning are the following:

- (1) A detailed evaluation of *Static* and *Dynamic* model selection strategies in a production environment.
- (2) An approach to enhance these strategies by incorporating energy usage metrics, significantly lowering energy consumption.

These contributions demonstrate that model selection strategies not only significantly reduce resource consumption but also have the potential to maintain or increase the accuracy of the AI system.

Moreover, these findings highlight the practicality and necessity of integrating Green AI principles into AI development, working towards more sustainable and efficient AI applications in the industry. The replication package for this study is available at the following DOI link: <https://doi.org/10.6084/m9.figshare.25481269>.

2 BACKGROUND

In this section, we highlight the concept of ensemble learning, followed by a detailed study of the use case and the infrastructure of the AI system.

2.1 Ensemble Learning

Ensemble learning is a strategic approach that combines several individual and diverse models to achieve better generalization performance [48]. The strength of ensemble learning lies in its diversity; a combination of models can provide a more robust and accurate output than any single model can [11, 18]. Popular applications of ensemble learning are speech recognition [14, 27] and classification [35, 43], image classification [23, 26, 45] and forecasting [5, 29, 37, 40].

A well-known method within ensemble learning is bagging (Bootstrap Aggregating), which enhances the stability and accuracy of machine learning algorithms by training multiple learners on various subsets of the original dataset and then aggregating their predictions to form a final decision [4].

At the core of our research are *model selection strategies* for ensembles. Model selection entails selecting a subset of models from an ensemble of models to optimize for a particular performance metric [6]. We use *model selection* within ensembles to reduce resource consumption in alignment with Green AI principles.

2.2 Use Case

To evaluate the impact of using *model selection strategies* for ensemble learning in a live production environment, we study the live AI system DocQMiner [12], a proprietary tool owned and developed by Deloitte NL. This system uses diverse machine learning (ML) models and NLP technologies to extract, process, and analyze data from textual documents. It is an information extraction tool widely used in the industry across many domains, with document set sizes ranging from 100 to over 100,000 documents. This paper focuses on the use case of extracting relevant properties from contracts.



Figure 1: Example of a document with relevant properties highlighted

After processing a document, the AI system makes predictions for predefined key properties. These properties are chosen by the user according to their use case, as shown in an example in Figure 1 where properties such as *Title*, *Parties*, and *Goal* in a contract are highlighted. The tool allows users to input a contract and delivers the predictions for the properties of interest. This design makes contract review more efficient, especially for complex documents [33].

DocQMiner employs diverse (pre-)trained models, as shown in Figure 2, to compile a ranked top 5 of textual predictions for every queried property. The AI system deploys an ensemble framework that utilizes the bagging approach. Multiple models that are independently trained on randomly generated subsets of data are combined to produce predictions.

For every property, each model in the ensemble produces a set of predictions aggregated to form the set of the five final predictions. The user selects the correct prediction, or, in case the desired answer is not part of the prediction, the user manually selects the answer from the document. This human-in-the-loop workflow ensures the validity of the processed data.

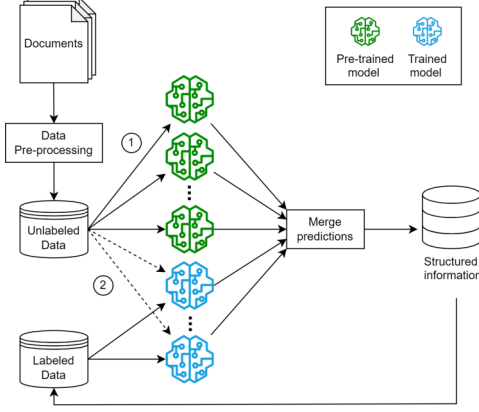


Figure 2: Workflow of DocQMiner. (1) Initially, documents are processed using the pre-trained models. After processing an initial set of documents, models are trained using the processed data. (2) After training models, documents are processed using the pre-trained and trained models.

The system subjects input documents to a pre-processing step to ensure that the text within is prepared for subsequent analysis by its models. Out of the box, the ensemble contains a set of pre-trained models to make predictions for the properties, indicated by the (1) in Figure 2. After the system has processed a set of documents within a domain, it can train an additional set of models using the in-domain data from those documents. Subsequently, the system processes documents using both the pre-trained models and those trained on the processed documents, as depicted by (2) in Figure 2.

2.3 Resource Consumption

This paper analyses the energy consumption of the AI system in a production environment. Therefore, we leverage the existing monitoring tool within the system, Datadog [10], which offers an extensive set of features for extracting metrics from the system. This cloud monitoring service provides a framework for tracking and analyzing energy use metrics within our AI deployments, as noted in industry literature [31].

DocQMiner uses a substantial amount of CPU when processing a single document. Considering DocQMiner processes document sets ranging in size from 100 to over 100,000 documents per instance, there are environmental and financial implications. This energy usage escalates operational costs and enlarges the carbon footprint of using DocQMiner, challenging the sustainable principles of Deloitte NL [13].

3 MODEL SELECTION

This study aims to analyze the impact of implementing model selection strategies in a live AI system that uses an ensemble of models on its energy usage and accuracy. In this section, we highlight how two selection strategies, *Static* [30] and *Dynamic* [9], can be used for this, and we describe our approach to using them to reduce resource consumption even more efficiently, *Energy-Aware selection*.

3.1 Model Selection for Energy Efficiency

In pursuit of sustainable AI practices, our study assesses existing *Static* [30] and *Dynamic* [9] model selection strategies to reduce energy consumption in a live AI system. These strategies are used to reduce the number of models or improve the efficiency of model usage used during inference, which is a significant determinant of overall energy usage [28].

The *Static* strategy selects an optimal subset of models for general tasks across the domain, while the *Dynamic* strategy adapts model selection to the specifics of each task, aiming to conserve energy without compromising the system’s accuracy.

Our main contribution lies in the novel *Energy-Aware selection* approach, which enhances the standard *Static* and *Dynamic* strategies by integrating an energy-aware metric in their application. This metric informs the selection process, ensuring that only the most energy-efficient models are chosen for the task at hand.

3.2 Static Selection

Static selection involves choosing the optimal subset of models for an entire domain. This approach generalizes the unique characteristics of the domain and selects the optimal subset across all queried properties. The following equation shows the process of *Static* selection.

$$S_{a^*} = \operatorname{argmax}_{S_i \in S} F1(S_i) \quad (1)$$

where S is the set of all possible model subsets, $F1(S_i)$ is the F1 score of subset S_i on the training set, and S_{a^*} as the optimal model subset chosen for evaluation.

3.3 Dynamic Selection

Dynamic selection is based on the belief that a model subset might not be optimal for an entire domain, but more specifically for the properties within the domain [9]. Therefore, we take the optimal subset for every queried property within the domain. This approach could be beneficial as the selection is more optimized per property. The process of *Dynamic* selection is represented by Equation 2.

$$S_{a^*}(p_j) = \operatorname{argmax}_{S_i \in S} F1(S_i, p_j) \quad (2)$$

We note P as the set of all properties within the domain. $F1(S_i, p_j)$ is the F1 score of subset S_i for property p_j on the training set. $S_{a^*}(p_j)$ is the optimal model subset for property p_j in P chosen for evaluation.

3.4 Energy-Aware Selection

Both selection strategies, *Static* and *Dynamic*, should reduce energy consumption because a subset of models is used instead of all. However, neither strategy considers how different models

compare in terms of energy efficiency. For instance, one model could slightly improve accuracy while costing a significantly larger amount of energy than another.

We propose an enhancement to the *Static* and *Dynamic* approach that discounts accuracy with resource consumption in the selection of models, *Energy-Aware selection*. We use the GreenQuotientIndex (GQI) [20] to factor the trade-off between accuracy and electricity usage. We add GQI to both *Static* and *Dynamic* versions, and compute it as follows:

$$GQI_{static} = \beta \times \frac{F1(S_i)^\alpha}{\log_{10}(C(S_i))} \quad (3)$$

$$GQI_{dynamic} = \beta \times \frac{F1(S_i, p_j)^\alpha}{\log_{10}(C(S_i))} \quad (4)$$

This metric evaluates the trade-off between accuracy and power consumption, where α and β are constants used to scale the GQI. The power consumption ($C(S_i)$) can vary significantly across different models, and therefore the logarithm of the power consumption is taken. Not all accuracy points ($F1(S_i)$ for GQI_{static} and $F1(S_i, p_j)$ for $GQI_{dynamic}$) have the same weight, as it is much easier to get from 0.4 to 0.5 than it is to go from 0.8 to 0.9. Therefore, the power (constant α) of the accuracy reflects the difference in difficulty.

Through the previously mentioned monitoring tool, Datadog [10], there is no availability for power consumption. We, therefore, use CPU usage as a proxy to discount the accuracy. We use this way of discounting the accuracy scores for both the *Dynamic* and the *Static* selection approaches highlighted above.

4 METHODOLOGY

This section outlines the methodology and evaluation of our *model selection strategies*. We start with data collection and analysis, followed by the experimental setup. We continue with the performance evaluation metrics. Lastly, we explore the energy consumption of the models during inference.

4.1 Data

The datasets for our experiments are selected based on a set of criteria. Each document within the dataset must contain contracts with a complete text, its associated properties, and annotated responses corresponding to these properties. We employ open-source datasets for our study, enhancing the reproducibility of our results. The CUAD [22] dataset and a dataset from the work of Leivaditi et al. [25] are used in this study, both of which were designed to optimize contract review processes and improve the effectiveness of information extraction.

4.1.1 CUAD. The Contract Understanding Atticus Dataset [22] (CUAD) is an extensive corpus tailored for commercial legal contract analysis. It comprises over 13,000 labels from 510 contracts divided into 25 contract types. Each document contains one or more of 41 distinct properties. The dataset presents a challenging research benchmark that can be used to enhance deep learning models' performance for contract analysis/understanding [7].

4.1.2 Lease Contracts. Leivaditi et al. [25] introduced a specialized benchmark dataset focused on lease agreement documents. These documents were sourced from a publicly available dataset by the U.S. Securities and Exchange Commission (SEC, 2020). The

dataset concentrates on extracting specific properties, including information about the lessor and details of the leased space.

Table 1: Relevant characteristics of CUAD and Lease Agreement dataset

Dataset	CUAD [22]	Lease Agreement [25]
#Documents	510	123
#TypesOfDocuments	25	1
#WordsPerDocument	7861	8053
#AnnotatedProperties	41	12
#TotalQueries	20,910	1476
#MissingAnnotations	13,959	494
MissingAnnotation (%)	66.77	33.47

4.1.3 Analysis of Datasets. We compare and analyze the datasets employed to understand the characteristics of the datasets and ensure the validity of the results. Table 1 shows the comparison of relevant characteristics. The CUAD and Lease Agreement datasets exhibit similar document lengths, with an average of approximately 8,000 words per document, as indicated by the *#WordsPerDocument* metric. Consequently, we will not regard document length as influencing our results.

The CUAD dataset encompasses a diverse collection of 25 types of contracts, reflecting a wide range of legal agreements. In contrast, the Lease Agreement dataset focuses solely on a single type of contract, specifically lease agreements. With the CUAD dataset's variety, the model selection process must account for the nuances and intricacies inherent in different types of contracts. Conversely, the Lease Agreement dataset's singular focus may allow for more specialized model tuning and optimization tailored specifically to lease agreements.

Additionally, there is a notable difference in the number of annotated properties between the two datasets. With *#MissingAnnotations*, we indicate the number of properties that were queried but did not have a ground truth in the document, thus missing an annotation. Specifically, 66.77% of the queried properties in the CUAD dataset lack annotations, in contrast to the Lease Agreement dataset, which is missing annotations for only 33.47% of its properties. Despite the significant number of missing annotations, we use these datasets to cover a broader array of use cases, including documents where the answer might not always exist. The varied percentage of missing annotations allows us to cover more ground in our results.

4.2 Experimental Setup

To integrate models tailored for specific domains, we select a subset of documents from the CUAD [22] and Lease Agreement [25] datasets that reflect the entire dataset. We process, annotate, and train models on these documents. After training, we run an evaluation on a held-out test set.

For both approaches, we run inference using the complete model ensemble on the documents from the validation set. We then evaluate the performance of the entire ensemble and each subset of models within the ensemble. The subset that demonstrates the highest performance, as determined by F1 scores from the training set, is selected for further testing.

This optimally performing subset is applied to the test set documents to evaluate its effectiveness on new data. To ensure the validity of the results, we create the validation and test datasets through 5-fold cross-validation.

As implemented in the tool, the predictions made by the entire model ensemble establish the performance evaluation baseline.

4.3 Performance Evaluation

When paired with the annotations, the ensemble’s predictions can result in various scenarios: True Positive, where the prediction aligns with the annotation; False Positive, where a prediction exists but fails to match the annotation; and False Negative, where a prediction exists but no corresponding annotation for the property exists.

To evaluate the performance of the strategies and compare it to the baseline situation of using all the models, we use the F1 score at k . F1 score is a balanced metric of Recall and Precision [36].

Precision: Precision [41] is the percentage of all predictions that match an annotation for the top k predictions:

$$\text{Precision@}k = \frac{TP}{TP + FP} \text{ for the top } k \text{ predictions} \quad (5)$$

Recall: Recall [41] is the percentage of correctly predicted annotations for the top k predictions:

$$\text{Recall@}k = \frac{TP}{TP + FN} \text{ for the top } k \text{ predictions} \quad (6)$$

Precision discounts for the number of False Positives and Recall for the number of False Negatives. Increasing Precision typically reduces Recall, and vice versa: increasing Recall decreases Precision. This trade-off shows the importance of prioritizing one over the other based on specific objectives.

F1 score: F1 score provides a balanced measure of Precision and Recall for the top k predictions. To ensure the robustness and reliability of our research, we select a subset of models based on the F1 score.

$$F1@k = 2 \cdot \frac{\text{Precision@}k \cdot \text{Recall@}k}{\text{Precision@}k + \text{Recall@}k} \quad (7)$$

DocQMiner [12] only shows the top 5 predictions produced by the set of models. Consequently, we evaluate the results based on k set to five. F1@5 serves as our primary criterion for selecting among different strategies, while we report Precision@5 and Recall@5 to provide a detailed view of the factors contributing to the F1@5 score.

4.4 Resource Consumption during Inference

To find the most energy-efficient model subset, we need a measurement of the models’ consumption during inference. DocQMiner is a live-production environment; therefore, we want to measure the system’s resource consumption live. Due to this limitation, we focused on tracking the Central Processing Unit (CPU) usage. With the use of Datadog [10], we can record the CPU usage per second for each process, which allows us to isolate the CPU usage per second, specifically during the inference phase of the models. By taking the cumulative sum of the CPU usage per second over the total duration of the process, we obtain the CPU seconds per process [3]:

$$\text{CPU seconds} = \sum_{i=1}^n u_i \quad (8)$$

n is the total number of seconds for the process, and u_i is the CPU usage per second during i . Datadog [10] performs this sum calculation.

This setup shows a clear view of each model’s performance, given that any other model does not influence the CPU utilization of one model. We gather the data for processing each document and each model for a set of intervals of the amount of queried properties.

Given that these measurements are taken in a live production environment, we designed our approach to yield results that closely represent the most probable outcomes. We acknowledge the inherent variability of a production environment, so we plot the measurements’ outcomes to account for the variance in results. To ensure the validity of the collected data, we repeat the measurement per number of queried properties 30 times.

5 RESULTS

Our evaluation of selection strategies across two datasets—CUAD and Lease Agreement—reveals significant differences in accuracy metrics, including Precision at 5 (P@5), Recall at 5 (R@5), F1 score, and the number of correct and incorrect predictions made by the models involved (see Table 2). *Full Ensemble* reflects the baseline of using all the models for every property within all documents.

5.1 Experimental Context

Legal information extraction. Extracting relevant information from legal documents presents significant challenges due to their complex nature. These texts often require identifying specific details within extensive documents. This task notably differs from more straightforward tasks like classification, where an F1 score below 0.6 might be deemed insignificant. In the context of legal text analysis, the F1 scores are typically lower, reflecting the intricate nature of the work involved.

These lower scores are supported in research conducted by Savelka et al. [38] using GPT-4 for property extraction from complex legal documents, which reported a Precision of 0.63, a Recall of 0.46, and an F1 score of 0.53. Despite being from a different dataset, these results show that relevant information extraction from legal documents is not a trivial task, and the results from the *Full Ensemble* should be interpreted as such.

Recall-oriented system. DocQMiner [12] is developed prioritizing Recall as the key metric, which aligns with its human-in-the-loop design. This design allows users to choose the best answer from the predictions provided. Consequently, the model development concentrates on accurately identifying relevant properties rather than minimizing the prediction of incorrect properties.

5.2 Full Ensemble

For CUAD, the *Full Ensemble* strategy achieved a P@5 of 0.0871, R@5 of 0.5338, and F1 of 0.1495, making 490 correct and 5928 incorrect predictions, with 100% CPU usage. For Lease Agreement, *Full Ensemble* had a P@5 of 0.1386, R@5 of 0.4448, and F1 of 0.2203, with 90 correct and 584 incorrect predictions, at 100% CPU usage.

5.3 Static Strategy

For CUAD, the *Static Original* variant, as defined in Section 3.2, achieved a P@5 of 0.6457, R@5 of 0.4872, and F1 of 0.5544, with

Table 2: Results of *Full Ensemble*, *Static* and *Dynamic* strategy in Precision@5, Recall@5, F1@5, Number of Correct Prediction, Number of Incorrect predictions and Consumption in percentage for CUAD and Lease Agreement dataset

Strategy		CUAD [22]						Lease Agreement [25]					
		P@5	R@5	F1@5	# Correct predictions	# Incorrect predictions	Relative CPU Usage	P@5	R@5	F1@5	# Correct predictions	# Incorrect predictions	Relative CPU Usage
<i>Full Ensemble</i>		0.0871	0.5338	0.1495	490	5928	100%	0.1386	0.4448	0.2203	90	584	100%
Static	Original	0.6457	0.4872	0.5544	408	265	60.39%	0.1763	0.3677	0.2545	77	376	64.62%
	Energy-Aware	0.6370	0.4652	0.5366	394	259	26.52%	0.1817	0.2404	0.2054	57	278	0.82%
Dynamic	Original	0.6118	0.5448	0.5750	446	322	71.28%	0.1852	0.4732	0.2652	93	427	81.22%
	Energy-Aware	0.6099	0.5471	0.5756	446	326	67.11%	0.2106	0.3415	0.2588	70	282	47.43%

408 correct and 265 incorrect predictions, consuming 60.39% CPU compared to the *Full Ensemble*. For Lease Agreement, *Static Original* posted a P@5 of 0.1763, R@5 of 0.3677, and F1 of 0.2545, with 77 correct and 376 incorrect predictions, at 64.62% CPU usage.

The *Static Energy-Aware* variation, as defined in Section 3.4, for CUAD showed a slight decrease in performance to an F1 of 0.5366, P@5 of 0.6370, R@5 of 0.4652, with 394 correct and 259 incorrect predictions, and reduced energy consumption to 26.52% CPU. For the Lease Agreement dataset, it managed an F1 of 0.2054, P@5 of 0.1817, and R@5 of 0.2404, with 57 correct and 278 incorrect predictions, drastically cutting CPU use to 0.82%.

Compared to the *Full Ensemble* baseline, overall Recall has slightly declined; however, overall Precision has increased, especially for the CUAD dataset. These results show that the *Static* strategy can correctly identify many properties while reducing the ‘noise’ of incorrect predictions.

5.4 Dynamic Strategy

For CUAD, the *Dynamic Original* strategy, as defined in Section 3.3, resulted in an F1 score of 0.5750, a P@5 of 0.6118, an R@5 of 0.5448, 446 correct and 322 incorrect predictions, at 71.28% CPU consumption. For Lease Agreement, it achieved an F1 score of 0.2652, a P@5 of 0.1852, an R@5 of 0.4732, 93 correct and 427 incorrect predictions, with 81.22% CPU usage compared to the *Full Ensemble*.

The *Dynamic Energy-Aware*, defined in Section 3.4, showed for CUAD an F1 of 0.5756, P@5 of 0.6099, and R@5 of 0.5471, with 446 correct and 326 incorrect predictions, lowering CPU usage to 67.11%. For the Lease Agreement dataset, the F1 was 0.2588, P@5 of 0.2106, R@5 of 0.3415, with 70 correct and 282 incorrect predictions, reducing CPU consumption to 47.43%.

Overall, the *Dynamic Original* strategy outperforms the baseline on both Precision and Recall. The increase in Recall is notable, considering DocQMiner is tailored to Recall. The increase in Recall is suspected to be due to how the ensemble ‘dilutes’ the predictions, and the *Dynamic* strategy specializes in specific properties.

5.5 Strategy Selection

Identifying an optimal strategy for an ensemble of models hinges on a few considerations. The evaluation of the *Full Ensemble*, *Static*, and *Dynamic* strategies across the CUAD and Lease Agreement datasets provides valuable insights into their respective strengths and weaknesses.

5.5.1 Precision - Recall Trade-off.

For Precision. The *Static Original* strategy stands out in the CUAD dataset with a Precision (P@5) of 0.6457 and an F1 score of 0.5544, significantly reducing the noise of incorrect predictions. Similarly, the Lease Agreement dataset performs with a Precision of 0.1763, compared to 0.1386 Precision for the *Full Ensemble*. These numbers suggest that the *Static Original* strategy offers a solution when minimising false positives is the goal.

For Recall. The *Dynamic Original* strategy shines by delivering a Recall (R@5) of 0.5448 for CUAD and 0.4732 for Lease Agreement, coupled with the highest F1 scores (0.5750 and 0.2652). This strategy ensures that more relevant properties are captured, making it ideal when maximising the number of accurately identified properties, which is the goal.

5.5.2 Resource efficiency. In a resource constraint environment, their performance and efficiency balance should inform the choice between the *Static Energy-Aware* and *Dynamic Energy-Aware* strategies.

Extreme Efficiency. The *Static Energy-Aware* strategy is unparalleled, especially evident in the Lease Agreement dataset, with CPU usage reduced to nearly 1%. This strategy is suitable for projects where every bit of computational resource saved makes an impact, even at the expense of some accuracy.

Balanced Approach. The *Dynamic Energy-Aware* strategy, while not as resource-efficient as its *Static* counterpart, offers a balance between accuracy and resource usage. This balance makes it ideal for scenarios where a moderate resource reduction is acceptable if it means retaining a higher level of accuracy.

5.5.3 Specificity of Properties. If the task involves identifying particular properties within legal documents, the specialisation afforded by the *Dynamic* strategies might yield better results. The *Dynamic* strategies are fine-tuned to identify specific properties more effectively, possibly at the cost of broader applicability that *Static* strategies can offer.

However, the success of the *Dynamic* approach hinges on the availability of sufficient information about the specific properties within the test set to identify the optimal subset accurately. In cases where such specific information is unavailable, the *Static* strategy may be the more suitable option, balancing the need for broader coverage with the available data.

5.6 Impact of Selection Strategies

Our results indicate that implementing model selection strategies can significantly impact both the accuracy and resource consumption of ensemble learning systems. The evaluation of the *Static* strategy suggests a notable improvement in Precision, reducing the number of incorrect predictions. In comparison, the *Dynamic* strategy excels in Recall, effectively retrieving more relevant instances while achieving a reduction in energy consumption, making it ideal for applications where capturing as much relevant information as possible is critical.

Furthermore, by enhancing these strategies with the cost-inclusive GreenQuotientIndex (GQI) [20], we have demonstrated a method for reducing the models' energy consumption without substantially sacrificing accuracy. For instance, the *Static Energy-Aware* strategy decreases the average energy usage from approximately 62% to 14% compared to the *Full Ensemble*, highlighting the potential for significant energy savings. In the case of the *Dynamic Energy-Aware* strategy, we observed an average reduction from around 76% to 57%, showing the effectiveness of this approach in balancing performance with energy efficiency.

These findings confirm that model selection strategies, in their original form and particularly when augmented with an energy consumption metric, can combine the objectives of maintaining high accuracy while reducing the energy demands of AI systems in a production environment.

6 DISCUSSION

This study presents several insights for strategic model selection for ensemble learning in live AI systems, particularly in performance optimization and computational efficiency.

6.1 (Green) AI

For AI practitioners, the findings emphasize the importance and viability of balancing accuracy with computational cost. In the current context, where computational efficiency is both an economic and environmental concern, the study's insights suggest that achieving this balance—maintaining high levels of accuracy while being mindful of the computational resources consumed—is essential for sustainable AI development [39].

These insights provide practical industry benefits, particularly in strategizing AI developments. The potential for resource-efficient model selection without significantly compromising performance paves the way for greener AI solutions. Such solutions are particularly valuable in resource-intensive applications, contributing to efforts to reduce AI technologies' environmental footprint.

6.2 Future of Ensembles

With the rising popularity of large language models (LLMs), one might wonder whether "old-school" ensembles are still the way to go. Research has shown that for most NLP tasks considered state-of-the-art fine-tuned models like TULRv6 generally outperform LLMs by a considerable margin, especially in languages other than English [1]. An LLM may not provide superior solutions for the specific task of legal information extraction.

In addition to accuracy, we prioritize efficiency as a crucial metric. A straightforward auto-regressive model like BERT_{Large}, which consists of approximately 340 million parameters [15], requires four

days of training on 16 TPU processors [16]. In contrast, GPT-4 is rumored to have 1.7 trillion parameters and demands 90-100 days of training on 25,000 GPU processors [17]. Although OpenAI has not officially disclosed the energy consumption of these models, it is reasonable to assume that their resource usage is significant. From the perspective of energy consumption, ensembles of fine-tuned models are preferable to large language models.

6.3 Monitoring

Monitoring is critical to ensuring the sustainability of AI systems [21, 47]. Practitioners need to be able to monitor their applications to understand their energy consumption and environmental impact. Tools and interfaces that facilitate the development of AI systems should also incorporate features for monitoring energy usage. Monitoring tools would enable practitioners to create eco-friendly systems without significant effort, addressing the gap in making sustainable AI more accessible and practical.

Moreover, by identifying the major energy consumers within live AI systems, practitioners can focus on reducing energy consumption in the most impactful areas, leading to more efficient and sustainable AI solutions.

6.4 Practical Implications

Balancing theoretical performance with practical considerations is vital when choosing the best model for a task. The *Static* strategy is typically more straightforward to deploy and compatible with a broader range of infrastructures, making it a practical option for many setups. The *Dynamic* strategy, while potentially more effective for specific tasks, might demand a more complex infrastructure setup. Therefore, the decision should consider the desired accuracy, efficiency, and practicality of integrating and maintaining the strategy into existing systems.

Companies should prioritize sustainable AI development as the impact of AI on the workforce continues to grow, accompanied by significant financial and environmental costs. Although the Green AI field is expanding, its connection to the industry remains insufficient [44]. This research demonstrates that making AI systems in production more sustainable is feasible. Continued efforts in this direction can further convince companies of the practicality and benefits of adopting sustainable AI practices.

6.5 Limitations

While this study provides valuable insights, it also has limitations. The evaluation was conducted on two specific datasets within the legal domain, which may limit the generalizability of the findings to other types of documents or domains. Future research could explore the applicability of these model selection strategies across a broader range of datasets and domains, not only within information extraction but also in other areas beyond the scope of NLP.

Additionally, our approach to discounting model selection based on resource consumption relies on CPU seconds as a proxy for energy consumption. Future studies could incorporate more direct metrics of energy consumption, such as power usage in kWh, to assess the environmental impact and reduction in the carbon footprint of different model selection strategies. This limitation aligns with the continued need for accurate and easy-to-use monitoring tools for AI systems [21, 47] highlighted in Section 6.3.

Lastly, our study does not account for the energy costs associated with training or fine-tuning models for specific domains within the ensemble. Future research should investigate whether the accuracy improvements from domain-specific training, see blue models in Figure 2, justify these increased energy expenditures. This analysis could provide deeper insights into the efficiency and effectiveness of employing specialized models within ensemble systems.

7 RELATED WORK

Green AI, the intersection of energy efficiency and model accuracy in Artificial Intelligence (AI), has sparked a growing interest amongst researchers [39, 42]. Our work focuses on reducing resource consumption in ensemble learning using model selection strategies in a live production environment. To our knowledge, there is no other work combining all elements. Below, we highlight the most significant contributions in model selection strategies and Green AI.

Zhou et al. [49] pivoted model selection in ensemble learning. Their work introduced GASEN, an approach that begins by allocating initial weights to neural networks and then employs a genetic algorithm to refine these weights. The optimized weights are instrumental in selecting the most effective subset of the ensemble. GASEN proves that incorporating the entire ensemble may not always be the optimal strategy, demonstrating the potential of selecting a subset of models.

Li et al. [28] present a novel approach, IRENE, to ensemble learning that focuses on balancing performance and computational cost for inference of ensemble learning. Using a learnable selector, base models, and implementing early halting in a sequential model setup, IRENE reduces inference costs by up to 56% while maintaining comparable performance to complete ensembles. IRENE presents a highly effective strategy for ensemble learning in contexts where sequential processing is feasible. However, the specific use case addressed in this paper necessitates a parallel approach and, as such, does not accommodate the sequential processing model that IRENE requires.

David et al. [11] propose an ensemble learning approach based on the consensus of multiple models for class prediction. It operates under the assumption that once multiple models predict a class, it is likely the correct one. This approach significantly reduces computational costs by about 50% while maintaining accuracy. Despite its effectiveness in classification tasks, this approach is not directly applicable to our work, as our ensemble is geared towards information extraction rather than classification, requiring a different methodological framework.

Kotary et al. [24] introduce a framework that combines machine learning and combinatorial optimization for differentiable model selection in ensemble learning. Their method of storing data for predicting optimal subsets in a neural network contrasts with our approach, which involves a more straightforward collection and selection of models within an ensemble for task-specific optimization. Additionally, as with the method from David et al. highlighted above, the success achieved on classification tasks cannot immediately be transferred to information extraction.

Cordeira et al. [9] present a new approach called Post-Selection Dynamic Ensemble Selection (PS-DES). PS-DES evaluates ensembles selected by different DES techniques using different metrics to

determine the best ensemble for each query instance. Their work introduces static and dynamic model selection based on evaluating a preliminary set of results. The relevance of these selection methods to our study stems from their design goal of integration into AI pipelines, which aligns with our focus on applying these strategies to a real-life use case. Consequently, we adopt these methods to reduce the number of models during inference in ensembles.

Within this landscape of optimizing ensemble learning, the work of Gowda et al. [20] introduces a critical consideration for assessing model efficiency. They propose a GreenQuotientIndex (GQI) metric that penalizes high electricity consumption while considering accuracy. The authors conduct a comprehensive study on the electricity consumption of different deep learning models, highlighting the often overlooked trade-off between accuracy and energy efficiency. Our work takes this concept further by demonstrating the practical application of the GQI in real-time production environments that rely on model ensembles.

8 CONCLUSION

In response to the rising energy demands of AI-powered software (AIware), particularly those employing ensemble learning [28], our empirical study investigates model selection strategies to optimize both accuracy and energy efficiency.

Our research introduces and evaluates two model selection strategies, *Static* and *Dynamic*, aimed at optimizing the performance of ensemble learning systems while minimizing their energy usage.

By evaluating the *Static* and *Dynamic* strategies across the CUAD and Lease Agreement datasets, we have highlighted the adaptability and potential of these approaches to meet diverse needs. Our results reveal that the *Static* strategy improves the F1 score beyond the baseline, reducing average energy usage from 100% from the full ensemble to 62%. The *Dynamic* strategy further enhances F1 scores, while using on average 76% compared to 100% of the full ensemble.

Additionally, we propose an approach that discounts accuracy with resource consumption, the *Energy-Aware* approach, showing potential for further reducing energy usage without significantly impacting accuracy. This method further decreased the average energy usage of the *Static* strategy from approximately 62% to about 14%, and for the optimal *Dynamic* strategy, from around 76% to roughly 57%.

Our findings, especially the successful application of the *Energy-Aware* approach, align with the principles of Green AI [39, 42], advocating for sustainable AI practices that maintain high-performance standards.

This field study on DocQMiner, an AI system actively used in the industry, highlights our research's real-world applicability and significance in advancing sustainable, efficient AI technologies for live production environments.

These insights provide a valuable perspective for the industry on developing AI in a resource-conscious yet effective manner. They highlight the feasibility of using model selection strategies to balance accuracy and computational efficiency, demonstrating the crucial need for adopting strategies that account for accuracy and environmental impact for ensuring sustainable development as AI progresses.

REFERENCES

- [1] Kabir Ahuja, Harshita Diddee, Rishav Hada, Millicent Ochieng, Krithika Ramesh, Prachi Jain, Akshay Nambi, Tanuja Ganu, Sameer Segal, Maxamed Axmed, et al. 2023. Mega: Multilingual evaluation of generative ai. *arXiv preprint arXiv:2303.12528* (2023).
- [2] Lasse F Wolff Anthony, Benjamin Kanding, and Raghavendra Selvan. 2020. Carbontracker: Tracking and predicting the carbon footprint of training deep learning models. *arXiv preprint arXiv:2007.03051* (2020).
- [3] James F Brady. 2005. Virtualization and cpu wait times in a linux guest environment. *J. Comp. Resource Mgmt* 116 (2005), 3–8.
- [4] Leo Breiman. 1996. Bagging predictors. *Machine learning* 24 (1996), 123–140.
- [5] Salvatore Carta, Andrea Corrigan, Anselmo Ferreira, Alessandro Sebastian Podda, and Diego Reforgiato Recupero. 2021. A multi-layer and multi-ensemble stock trader using deep learning and deep reinforcement learning. *Applied Intelligence* 51 (2021), 889–905.
- [6] Rich Caruana, Alexandru Niculescu-Mizil, Geoff Crew, and Alex Ksikes. 2004. Ensemble selection from libraries of models. In *Proceedings of the twenty-first international conference on Machine learning*. 18.
- [7] Yupeng Chang, Xu Wang, Jindong Wang, Yuan Wu, Kaijie Zhu, Hao Chen, Linyi Yang, Xiaoyuan Yi, Cunxiang Wang, Yidong Wang, et al. 2023. A survey on evaluation of large language models. *arXiv preprint arXiv:2307.03109* (2023).
- [8] Andrew A Chien, Liuzixuan Lin, Hai Nguyen, Varsha Rao, Tristan Sharma, and Rajini Wijayawardana. 2023. Reducing the Carbon Impact of Generative AI Inference (today and in 2035). In *Proceedings of the 2nd Workshop on Sustainable Computer Systems*. 1–7.
- [9] Paulo RG Cordeiro, George DC Cavalcanti, and Rafael MO Cruz. 2023. A post-selection algorithm for improving dynamic ensemble selection methods. *arXiv preprint arXiv:2309.14307* (2023).
- [10] DataDog. 2024. DataDog: Cloud Monitoring as a Service. <https://www.datadoghq.com/>. Accessed: 2024-02-08.
- [11] Nelly David and Nathan S Netanyahu. 2021. Adaptive consensus-based ensemble for improved deep learning inference cost. In *International Conference on Artificial Neural Networks*. Springer, 330–339.
- [12] Deloitte. 2022. DocQMiner - Deloitte's Solution for Document Processing and Information Extraction. <https://www2.deloitte.com/nl/nl/pages/risk/solutions/docqminer.html>. Accessed: 2024-02-11.
- [13] Deloitte Netherlands. 2023. Environmental Impacts - Deloitte Netherlands Annual Report 2022/2023. <https://annualreport.deloitte.nl/annual-report-20222023/nonfinancial-statements--2-environmental-impacts>. Accessed: 2024-03-27.
- [14] Li Deng, Gokhan Tur, Xiaodong He, and Dilek Hakkani-Tur. 2012. Use of kernel deep convex networks and end-to-end learning for spoken language understanding. In *2012 IEEE Spoken Language Technology Workshop (SLT)*. IEEE, 210–215.
- [15] Hugging Face. 2020. Transformers: Pretrained Models. https://huggingface.co/transformers/v2.4.0/pretrained_models.html. Accessed: 2024-03-28.
- [16] Hugging Face. 2021. BERT 101. <https://huggingface.co/blog/bert-101>. Accessed: 2024-03-28.
- [17] Danielle Franca. 2023. GPT4: All Details Leaked. <https://medium.com/@daniellefranca96/gpt4-all-details-leaked-48fa20f9a4a>. Accessed: 2024-03-28.
- [18] Mudasiir A Ganaie, Minghui Hu, AK Malik, M Tanveer, and PN Suganthan. 2022. Ensemble deep learning: A review. *Engineering Applications of Artificial Intelligence* 115 (2022), 105151.
- [19] Stefanos Georgiou, Maria Kechagia, Tushar Sharma, Federica Sarro, and Ying Zou. 2022. Green ai: Do deep learning frameworks have different costs? In *Proceedings of the 44th International Conference on Software Engineering*. 1082–1094.
- [20] Shreyank N Gowda, Xinyue Hao, Gen Li, Laura Sevilla-Lara, and Shashank Narayana Gowda. 2023. Watt For What: Rethinking Deep Learning's Energy-Performance Relationship. *arXiv preprint arXiv:2310.06522* (2023).
- [21] Peter Henderson, Jieru Hu, Joshua Romoff, Emma Brunskill, Dan Jurafsky, and Joelle Pineau. 2020. Towards the systematic reporting of the energy and carbon footprints of machine learning. *The Journal of Machine Learning Research* 21, 1 (2020), 10039–10081.
- [22] Dan Hendrycks, Collin Burns, Anya Chen, and Spencer Ball. 2021. CuaD: An expert-annotated nlp dataset for legal contract review. *arXiv preprint arXiv:2103.06268* (2021).
- [23] Gao Huang, Yixuan Li, Geoff Pleiss, Zhuang Liu, John E Hopcroft, and Kilian Q Weinberger. 2017. Snapshot ensembles: Train 1, get m for free. *arXiv preprint arXiv:1704.00109* (2017).
- [24] James Kotary, Vincenzo Di Vito, and Ferdinando Fioretto. 2023. Differentiable model selection for ensemble learning. In *International Joint Conference on Artificial Intelligence (IJCAI)*.
- [25] Spyretta Leivaditi, Julien Rossi, and Evangelos Kanoulas. 2020. A benchmark for lease contract review. *arXiv preprint arXiv:2010.10386* (2020).
- [26] Jichang Li, Si Wu, Cheng Liu, Zhiwen Yu, and Hau-San Wong. 2019. Semi-supervised deep coupled ensemble learning with classification landmark exploration. *IEEE Transactions on Image Processing* 29 (2019), 538–550.
- [27] Sheng Li, Xugang Lu, Shinsuke Sakai, Masato Mimura, and Tatsuya Kawahara. 2017. Semi-supervised ensemble DNN acoustic model training. In *2017 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 5270–5274.
- [28] Ziyue Li, Kan Ren, Yifan Yang, Xinyang Jiang, Yuqing Yang, and Dongsheng Li. 2023. Towards Inference Efficient Deep Ensemble Learning. *arXiv preprint arXiv:2301.12378* (2023).
- [29] Fan Liu, Feng Xu, and Sai Yang. 2017. A flood forecasting model based on deep learning algorithm via integrating stacked autoencoders with BP neural network. In *2017 IEEE third International conference on multimedia big data (BigMM)*. Ieee, 58–61.
- [30] Gang Luo. 2016. A review of automatic selection methods for machine learning algorithms and hyper-parameter values. *Network Modeling Analysis in Health Informatics and Bioinformatics* 5 (2016), 1–16.
- [31] Adam Mckaig and Tahia Khan. 2022. How the Metrics Backend Works at Datadog. (2022).
- [32] Dagmar Monett and Claudia Lemke. 2021. AI-ware: Bridging AI and Software Engineering for responsible and sustainable intelligent artefacts. *Managing Artificial Intelligence* (2021), 66–73.
- [33] Afra Nawar, Mohammed Rakib, Salma Abdul Hai, and Sanaulla Haq. 2022. An Open Source Contractual Language Understanding Application Using Machine Learning. In *Proceedings of the First Workshop on Language Technology and Resources for a Fair, Inclusive, and Safe Society within the 13th Language Resources and Evaluation Conference*. 42–50.
- [34] Deloitte Netherlands. 2024. Home Page. <https://www2.deloitte.com/nl/nl.html>. [Online; accessed March 6, 2024].
- [35] Hamid Palangi, Li Deng, and Rabab K Ward. 2014. Recurrent deep-stacking networks for sequence classification. In *2014 IEEE China Summit & International Conference on Signal and Information Processing (ChinaSIP)*. IEEE, 510–514.
- [36] David MW Powers. 2020. Evaluation: from precision, recall and F-measure to ROC, informedness, markedness and correlation. *arXiv preprint arXiv:2010.16061* (2020).
- [37] Xueheng Qiu, Le Zhang, Ye Ren, Ponnuthurai N Suganthan, and Gehan Amarathunga. 2014. Ensemble deep learning for regression and time series forecasting. In *2014 IEEE symposium on computational intelligence in ensemble learning (CIEL)*. IEEE, 1–6.
- [38] Jaromir Savelka, Kevin D Ashley, Morgan A Gray, Hannes Westermann, and Huihui Xu. 2023. Can GPT-4 Support Analysis of Textual Data in Tasks Requiring Highly Specialized Domain Expertise? *arXiv preprint arXiv:2306.13906* (2023).
- [39] Roy Schwartz, Jesse Dodge, Noah A Smith, and Oren Etzioni. 2020. Green ai. *Commun. ACM* 63, 12 (2020), 54–63.
- [40] Pardeep Singla, Manoj Duhan, and Sumit Saroha. 2022. An ensemble method to forecast 24-h ahead solar irradiance using wavelet decomposition and BiLSTM deep learning network. *Earth Science Informatics* 15, 1 (2022), 291–306.
- [41] Marina Sokolova and Guy Lapalme. 2009. A systematic analysis of performance measures for classification tasks. *Information processing & management* 45, 4 (2009), 427–437.
- [42] Emma Strubell, Ananya Ganesh, and Andrew McCallum. 2019. Energy and policy considerations for deep learning in NLP. *arXiv preprint arXiv:1906.02243* (2019).
- [43] Gokhan Tur, Li Deng, Dilek Hakkani-Tür, and Xiaodong He. 2012. Towards deeper understanding: Deep convex networks for semantic utterance classification. In *2012 IEEE international conference on acoustics, speech and signal processing (ICASSP)*. IEEE, 5045–5048.
- [44] Roberto Verdecchia, June Sallou, and Luís Cruz. 2023. A systematic review of Green AI. *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery* (2023), e1507.
- [45] Bin Wang, Bing Xue, and Mengjie Zhang. 2020. Particle swarm optimisation for evolving deep neural networks for image classification by evolving and stacking transferable blocks. In *2020 IEEE Congress on Evolutionary Computation (CEC)*. IEEE, 1–8.
- [46] Tim Whitaker and Darrell Whitley. 2022. Prune and tune ensembles: low-cost ensemble learning with sparse independent subnetworks. In *Proceedings of the AAAI Conference on Artificial Intelligence*, Vol. 36. 8638–8646.
- [47] Carole-Jean Wu, Ramya Raghavendra, Udit Gupta, Bilge Acun, Newsha Ardalani, Kiwan Maeng, Gloria Chang, Fiona Aga, Jinshi Huang, Charles Bai, et al. 2022. Sustainable ai: Environmental implications, challenges and opportunities. *Proceedings of Machine Learning and Systems* 4 (2022), 795–813.
- [48] Zhi-Hua Zhou. 2012. *Ensemble methods: foundations and algorithms*. CRC press.
- [49] Zhi-Hua Zhou, Jianxin Wu, and Wei Tang. 2002. Ensembling neural networks: many could be better than all. *Artificial intelligence* 137, 1-2 (2002), 239–263.

Received 2024-04-05; accepted 2024-05-04