



Review

Autophagy and machine learning: Unanswered questions

Ying Yang^{a,b,c,d,1}, Zhaoying Pan^{c,1}, Jianhui Sun^e, Joshua Welch^{c,d}, Daniel J. Klionsky^{a,b,*}^a Department of Molecular, Cellular, and Developmental Biology, University of Michigan, Ann Arbor, MI 48109, USA^b Life Sciences Institute, University of Michigan, Ann Arbor, MI 48109, USA^c Department of Electrical Engineering and Computer Science, University of Michigan, Ann Arbor, MI 48109, USA^d Department of Computational Medicine and Bioinformatics, University of Michigan, Ann Arbor, MI 48109, USA^e Department of Computer Science, University of Virginia, Charlottesville, VA 22903, USA

ARTICLE INFO

Keywords:

Artificial intelligence

Lysosome

Macroautophagy

Stress

ABSTRACT

Autophagy is a critical conserved cellular process in maintaining cellular homeostasis by clearing and recycling damaged organelles and intracellular components in lysosomes and vacuoles. Autophagy plays a vital role in cell survival, bioenergetic homeostasis, organism development, and cell death regulation. Malfunctions in autophagy are associated with various human diseases and health disorders, such as cancers and neurodegenerative diseases. Significant effort has been devoted to autophagy-related research in the context of genes, proteins, diagnosis, etc. In recent years, there has been a surge of studies utilizing state of the art machine learning (ML) tools to analyze and understand the roles of autophagy in various biological processes. We taxonomize ML techniques that are applicable in an autophagy context, comprehensively review existing efforts being taken in this direction, and outline principles to consider in a biomedical context. In recognition of recent groundbreaking advances in the deep-learning community, we discuss new opportunities in interdisciplinary collaborations and seek to engage autophagy and computer science researchers to promote autophagy research with joint efforts.

1. Introduction

We have published three articles as part of an “unanswered question” series of reviews, entitled “The molecular machinery of autophagy: unanswered questions” [1], “The regulation of autophagy: unanswered questions” [2], and “Autophagy and diseases: unanswered questions” [3], discussing fundamental concepts of the autophagy mechanism and unanswered aspects of its involvement in various human diseases. With the increasingly emerging large-scale sets of data, there is a significant need to integrate and utilize these data to provide novel insights to understand the complicated process of autophagy with ML tools. In this review, we provide a comprehensive view of the progress made in applying ML algorithms to autophagy research. In our discussion, we explore four key research categories: gene/protein identification, biomarker identification, drug discovery, and morphological analysis. Each category represents a distinct topic of autophagy research where ML approaches have shown significant promise. By exploring these diverse topics, we aim to highlight the versatility and potential of ML applications in advancing our understanding of autophagy across various research domains. We also provide questions and potential

directions that may bring new breakthroughs in our understanding and offer promising opportunities for interdisciplinary collaborations.

2. Research questions

In this review, we address and provide insights into the following research questions, catering to a diverse audience with an autophagy or ML background to promote future interdisciplinary research:

- (1) What machine learning models are currently utilized in autophagy research?
- (2) What kind of autophagy research has been explored with machine learning?
- (3) What will be the promising directions for interdisciplinary research of autophagy and machine learning?

To respond to the first question, an in-depth analysis is conducted in the subsection titled “Autophagy Machinery and Machine Learning Concepts.” The second research question is answered in the “Recent Advances” section with four distinct research topics. Last, the third question is discussed in the “Open Questions” section, providing potential directions for future research. We show an overview of our paper

* Corresponding author at: Department of Molecular, Cellular, and Developmental Biology, University of Michigan, Ann Arbor, MI, 48109, USA.

E-mail address: klionsky@umich.edu (D.J. Klionsky).

¹ These authors contributed equally to the work.

in Fig. 1.

3. Autophagy machinery and machine learning concepts

As an interdisciplinary review targeting an audience with potentially diverse backgrounds, we start with basic concepts of autophagy and ML. In this section, we first briefly review the fundamental concepts of autophagy and highlight its crucial role in diverse aspects of cellular physiology. We aim to bridge the gap between autophagy research and ML by introducing key models and evaluation metrics that have been frequently utilized in current autophagy research.

3.1. Autophagy concepts

Autophagy serves as a critical cellular process that balances sources of energy during key development stages and in response to nutrient stress. With its role in degrading and recycling misfolded or aggregated proteins, damaged organelles, and intracellular pathogens, autophagy—or autophagy dysfunction—is widely connected with the development of diseases, including cancer, metabolic disorders and neurodegenerative diseases [4].

Depending on the mechanism of how cargos are delivered to lysosomes for degradation, there are three primary types of autophagy, including macroautophagy, microautophagy, and chaperone-mediated autophagy (CMA) [5].

Macroautophagy, the most widely investigated form, holds a key position in the development of numerous diseases linked to autophagic irregularities/dysfunction. Its evolutionarily conserved process encompasses five stages: initiation, nucleation, expansion, fusion, and degradation and recycling (Fig. 2) [6]. Morphologically, macroautophagy is distinguished by the creation of autophagosomes, which are double-membrane structures facilitating the transport of damaged organelles and protein aggregates to lysosomes. The process of macroautophagy is characterized by the hierarchical involvement of ATG (autophagy related) proteins.

Macroautophagy can be categorized into two main types: selective autophagy and nonselective (also referred to as bulk) autophagy. Selective autophagy is distinguished by its high specificity in selecting and delivering specific cargo for degradation [7]. Conversely, nonselective autophagy is thought to randomly engulf cytoplasm and lacks specificity in the cargo it targets for degradation. Mitophagy, the selective autophagic degradation of mitochondria, is the most comprehensively studied form of selective autophagy; it holds a significant role in

safeguarding against numerous diseases, especially various neurodegenerative conditions. This process aids in the removal of damaged mitochondria and contributes to the regulation of mitochondrial quality control.

Microautophagy, which entails the direct translocation of cytosolic material into lysosomes via membrane invagination or protrusion, contributes significantly to cellular homeostasis [8]. While its involvement in disease is not as thoroughly understood, any dysregulation has the potential to affect cellular health.

CMA, a distinctive mechanism that selectively transports individual, unfolded substrate proteins into lysosomes without substantial membrane alterations, plays a pivotal role in maintaining cellular quality control [9]. CMA is particularly known as a mechanism for safeguarding neurons against damage in neurodegenerative diseases, including brain injury and the aggregation of toxic proteins.

Important research topics in autophagy encompass revealing new genes or molecules associated with autophagy, understanding the molecular machinery governing the process, identifying regulatory pathways, and elucidating autophagy's role in diverse biological contexts. Additionally, autophagy's involvement in human health and disease is a major focus. In this review, we refrain from specifying particular autophagic types; instead focusing on comprehensive research for autophagy integrated with ML tools. We identify several research topics that have applied ML tools and call for broader application of these tools to advance research on these topics.

3.2. Machine learning concepts

This subsection provides an overview of commonly used ML techniques in autophagy research. ML can be generally divided into supervised and unsupervised learning. Unsupervised learning deals with unlabeled data, where the model seeks to find patterns and relationships within the model without explicit guidance. Supervised learning involves training a model using labeled data, where the input data are paired with corresponding output labels [10]. The goal is to learn the mapping between inputs and outputs so that accurate predictions can be made when presented with unseen data. Classification and regression are two fundamental tasks in supervised learning. Classification involves categorizing data into predefined classes based on input features, and regression predicts continuous numeric values based on input features. Additionally, supervised learning can be divided into traditional models and deep-learning models. Traditional ML models such as decision trees and linear regression offer transparent decision-making progress and

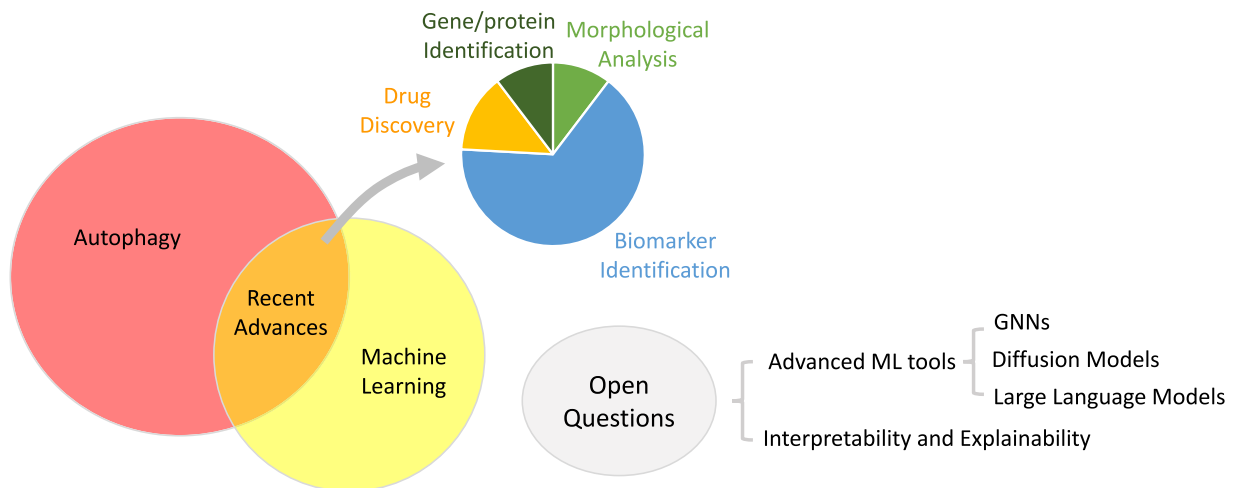


Fig. 1. An Overview of the Review. We explore the recent advance of autophagy research involving ML tools, categorized into four different tasks, including drug discovery, gene/protein identification, morphological analysis, and biomarker identification. We then address open questions, including the application of advanced ML tools as well as the interpretability and explainability of ML outcomes.

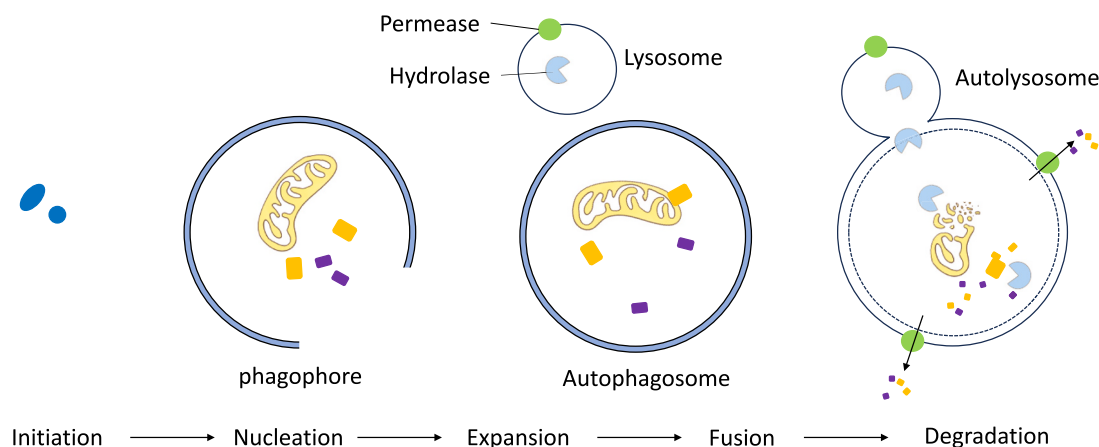


Fig. 2. Schematic representation of nonselective macroautophagy. The initiation of macroautophagy occurs in response to cellular stresses. A structure known as the phagophore emerges and expands to encapsulate unwanted organelles and macromolecules, forming a double-membrane structure referred to as the autophagosome. Subsequently, the outer membrane of the autophagosome merges with a lysosome, facilitating the transport of the cargo to the newly created autolysosome for degradation. The resultant breakdown products, such as amino acids, fatty acids, and nucleotides, are further released via permeases and then recycled and used for the synthesis of new cellular structures or as an energy source.

perform well with scarce data in a computationally efficient way compared with deep-learning models. Deep-learning models can learn complex patterns and relationships in data and outperform traditional ones on large-scale datasets.

Most of the learning tasks in autophagy research can be formulated as classification problems, where researchers aim to determine the functional types of specific genes or identifying biomarkers of particular diseases, both of which are essentially inferring ‘labels’ based on biological features. Therefore, we start with introducing models that are most frequently used in classification, including logistic regression (LR), LASSO Cox regression [11], k-nearest neighbors (KNN) [12,13], random forest (RF) [14,15], eXtreme Gradient Boosting (XGBoost) [16], support vector machine (SVM) [17] and SVM with recursive feature elimination (SVM-RFE) [18]. These models could be categorized into traditional ML models. While these models can be used beyond classification use cases, this review mainly focuses on how they can help solve classification problems arising in autophagy research. There is no one-size-fits-all approach to select the most appropriate model, which is entirely task-dependent, and it is highly recommended to conduct exploratory data analysis and experiment with a variety of models. Scientific discovery is usually extremely rigorous, thus requiring that the models’ results are robust, i.e., consistently reproducible by a variety of independently run models, and explainable, i.e., the decision logic of trained models makes scientific sense.

In addition, we briefly talk about clustering, an unsupervised learning approach that is able to recognize patterns in a label-scarce regimen and is widely used in the task of disease subtype identification. We further introduce several deep-learning models that have shown promise in current research, such as convolutional neural networks (CNNs) [19], recurrent neural networks (RNNs) [20], graph neural networks (GNNs) [21], biological text mining models such as Mol2vec [22] and BioBERT [23]. We also discuss the evaluation metrics of the model performance in the subsequent section. To our best knowledge, many of the groundbreaking innovations evolving in recent years, including diffusion models [24–27] and large language models (LLMs) [28–32] have yet to be applied in autophagy research. This gap presents exciting opportunities for researchers to explore how these powerful tools can be harnessed in an autophagy context. We include a short discussion of these models in the “Open Questions” section of this paper.

3.2.1. Logistic regression

LR is a straightforward model arising from binary classification tasks.

LR applies a logistic transformation of the input features to estimate the likelihood of the given sample belonging to one of the two classes. This model assumes a straightforward linearity between features and labels, which makes the model highly interpretable, while restricting the model’s representation capacity.

LR is a great example of this fundamental model complexity – interpretability tradeoff. Whereas more complex models are capable of inferring a more intricate relationship (beyond linearity) between features and labels, they lack transparency and obscure the logic behind predictions.

3.2.2. LASSO Cox Regression

LASSO Cox Regression integrates a Cox model, which originates from survival analysis where time to event is predicted based on features, and LASSO regression which essentially penalizes a model for including irrelevant features. LASSO Cox regression identifies and prioritizes the most relevant features associated with the event of interest, incorporating LASSO regularization to prevent overfitting and select important features. Careful tuning of regularization parameters is crucial for optimal performance.

3.2.3. K-nearest neighbors

KNN is a simple yet effective method in ML for classification tasks. Based on a given metric of distance, KNN predicts the label of a new data point by taking a majority vote among its K (which is a hyperparameter to pre-specify) nearest neighbors in the training dataset. Despite its simplicity for implementation, it may be computationally expensive with large datasets or features of high dimensions.

3.2.4. Random forest

RF, an ensemble learning method, builds upon a number of decision trees. Decision trees classify data by segmenting the dataset into smaller subsets through a sequence of rules. RF ensembles multiple base decision trees and aggregates the majority voting as the classification results. Compared to decision trees, RF reduces the risk of overfitting but sacrifices the interpretability.

3.2.5. XGBoost

XGBoost is a powerful ensemble learning method renowned for its high efficiency and effectiveness. In XGBoost, decision trees are built in parallel, and the errors made by previous trees are corrected iteratively by the successive trees to enhance the performance. Although XGBoost has been applied in a variety of tasks, it requires careful tuning of

parameters to achieve optimal performance.

3.2.6. Support vector machine

Data samples can be conceptualized as points in a high-dimensional space, and SVM aims to find an optimal hyperplane that best separates data into different classes. SVM can handle complex decision boundaries, is better in generalization, and can model a non-linear decision boundary, but its performance depends on proper hyperparameter tuning. SVM with Recursive Feature Elimination (SVM-RFE), built upon SVM, iteratively ranks features and eliminates least important features. SVM-RFE is widely used in the analysis of biological data; however, this feature selection process can be computationally expensive.

3.2.7. Clustering

Clustering is a fundamental technique in unsupervised learning. In contrast to supervised learning, unsupervised learning does not require labels for data; instead, it utilizes the data itself to reveal its patterns. With a prespecified number of clusters, data are clustered into groups in such a way that within-cluster data are much more similar than cross-cluster data. In practice, the selection of distance metrics, and number of clusters are key hyperparameters to obtain optimal performance.

3.2.8. Convolutional neural networks

CNNs are deep neural networks that are commonly used in structured data such as images. CNNs use a series of neural network layers to process images for different tasks, and the convolutional layers are a key design to leverage spatially correlated pixels to engineer task-specific features from images. Among various variants, U-Net [33] and ResNet [34] stand out as two of the most outstanding models. U-Net excels in image segmentation tasks, while ResNet is renowned for its deep architecture enabling advanced image classification.

Notably, deep neural networks including CNNs, include a significant amount of learnable parameters (from millions to billions), thus requiring much more data for training for optimal performance and to avoid overfitting.

3.2.9. Recurrent neural networks

RNNs are neural networks specially designed to process sequence data. The key feature of RNNs is their inherent ability to retain and utilize information from past inputs. As the model processes each element in a sequence, it incorporates relevant details from all preceding elements, allowing it to build and maintain a contextual understanding over time. However, RNNs often struggle to capture the dependency between two distant positions of a sequence, due to problems such as vanishing or exploding gradients. This has led to the development of more advanced variants like Long Short-Term Memory/LSTM [35] and Gated Recurrent Unit/GRU [36] networks, which are designed to better capture long-term dependencies.

3.2.10. Graph neural networks

GNNs are powerful tools for analyzing graph data, a data structure representing the relations (edges) between a set of entities (nodes). The basic logic to obtain node representation in GNNs is to aggregate and transform feature information from the immediate neighbors of that node, effectively allowing the graph structure to be learned. Graph Convolutional Networks/GCNs [37] and Graph Attention Networks/GATs [38] represent significant improvements in GNNs. Graph Convolutional Networks are suitable in analyzing uniform graph structures, while Graph Attention Networks are optimal for graphs with non-uniform influence among neighbor nodes. However, challenges arise with GNNs in handling large-scale dense graphs, when the number of neighbors (adjacent or k-hop) grow exponentially, and in ensuring robustness against irregular graph structures.

3.2.11. Natural language processing

NLP models are deep-learning methods designed to process text

sequences. Biological sequences e.g. molecule representations and RNA-seq resemble text sequences in many ways, and NLP approaches have demonstrated its applicability in the biological domain. Word2vec [39] is among the earliest successes to obtain dense vector representations (also known as embeddings) of words. The representations of words with similar semantic meanings are closer in the vector space. Inspired by Word2vec, Mol2vec [22] was proposed to learn the embeddings of compounds by aggregating the vectors of substructures of molecules. Large NLP models (also known as Large Language Models) can also be used as an external knowledge base because they are usually trained with an enormous text corpus and can comprehend much domain knowledge. For example, Bidirectional Encoder Representations from Transformers (BERT) [29] is a powerful language model for various downstream tasks in NLP, trained on a massive amount of text data to learn the contextual relationships of words in natural language. As an application of BERT, BioBERT [23] is a domain-specific language model trained on large-scale biomedical corpora, to help researchers understand complex biomedical texts.

3.2.12. Usage of ML techniques

In this section, we provide a brief discussion on the usage of the above ML techniques, including the commonly used preprocessing techniques, the summary of their applications, evaluation metrics, and possible biases.

Preprocessing data is a crucial step in machine learning pipelines to prepare the data for effective modeling and ensure optimal performance. Common preprocessing includes data filtering, normalization, feature encoding, and feature selection. Data filtering involves removing broken, noisy, or irrelevant data to improve data quality. Normalization scales the data to a standard range, ensuring that features with different scales have equal influence in the analysis. Feature encoding transforms data with various formats into numerical representations, making them suitable for use in machine learning algorithms. Feature selection involves selecting features from the original dataset to capture relevant information and improve model performance. However, it is important to note that data preprocessing techniques may vary depending on the characteristics of the data and the objectives of the task; the specific techniques employed should be carefully evaluated and tailored according to the data and task in practice.

In autophagy research, the utilization of ML techniques has been expanding, with a predominant focus on binary classification tasks using traditional ML models. These tasks involve categorizing instances into two classes, such as whether a gene is autophagy related or not. Models including LR, KNN, RF, XGBoost, SVM, and SVM-RFE are commonly applied for binary classification and can be extended to handle multi-class classification scenarios. As the field progresses, there is a growing interest in leveraging more advanced techniques in autophagy research. Models such as CNNs, RNNs, and GNNs offer the capability to capture complex patterns of image, sequence, or graph in autophagy-related data, making them suitable for feature encoding and classification tasks. Additionally, biological text mining models such as Mol2vec and BioBERT are increasingly utilized for analyzing or encoding unstructured biomedical text data. Moreover, feature selection plays a crucial role in identifying relevant features for model performance improvement. Models including RF, LASSO Cox regression, XGBoost, SVM with linear kernels, and SVM-RFE are commonly employed for feature selection, enabling researchers to identify key features associated with the ML predictions. As the field continues to evolve, it is essential to explore the full potential of these diverse ML models and techniques in addressing complex questions in autophagy research while being mindful of the specific characteristics and challenges in different scenarios.

Depending on each task, there are various metrics for performance evaluation. We first introduce the metrics for classification task, which is predominant in autophagy research. The confusion matrix serves as the foundation for many subsequent numeric metrics, dividing all model

predictions into four categories: true positives (TP), false positives (FP), true negatives (TN), and false negatives (FN). We show a visualization of a confusion matrix in Fig. 3(a). Based on this matrix, the following metrics emerge:

Accuracy is the proportion of correctly classified instances of both positive and negative instances, calculated as $(TP + TN) / (TP + FP + TN + FN)$. Accuracy serves as the default metric for most of the tasks, but it may not be the most appropriate choice in imbalanced data classification. Precision, formulated as $TP / (TP + FP)$, describes the ability of positive predictions, which is extremely useful when the impact of false positives is significant. Recall, also known as sensitivity, or true positive rate, calculated as $TP / (TP + FN)$, represents the model's capacity in all actual positive instances, which is important when failing to detect positive instances can lead to serious consequences. Specificity, also true negative rate, is defined as $TN / (TN + FP)$, which is also $1 - \text{false positive rate}$, and is helpful when incorrectly identifying negative instances as positive leads to critical repercussion.

The F1 Score, a harmonic mean of precision and recall, calculated as $2 * \text{precision} * \text{recall} / (\text{precision} + \text{recall})$, provides a balanced assessment of a model's performance, particularly in datasets with imbalanced class distributions. In classification tasks, we use a cut-off threshold to determine the predicted class membership with the likelihood outputted by the models. The receiver operating characteristic (ROC) curve assesses the model's classification ability by plotting the true positive rate against the false positive rate (calculated as $FP / [FP + TN]$) across various decision thresholds. The area under the curve (AUC) is the total area under the ROC curve, and a higher AUC value indicates better classification performance. In Fig. 3(b), we show examples of ROC curves, and provide the AUC value, 0.5, for a random classifier as an example.

The metrics for binary classification can be extended to evaluate the performance of multi-class classification on individual classes. Moreover, to better evaluate the overall performance for multi-class classification, micro-averaged metrics calculate performance across all classes by treating each instance equally, while macro-averaged metrics provide equal weight to each class, regardless of their prevalence.

Clustering, one of the techniques mentioned in the previous section, has been utilized in the current work as well. Instead of aiming to learn to cluster existing data and categorize new data correctly, the application of clustering has been limited to identifying the existence of potential clusters (subtypes). Therefore, the clustering can be validated by the interpretation of domain experts without explicit quantitative evaluation. If necessary, the Silhouette coefficient is a popular metric to evaluate the quality of clustering. In contrast, regression, another fundamental task in ML, was not utilized widely in current ML

applications in autophagy research. Compared to classification that categorizes instances into discrete labels, regression models the relationship between variables and predicts continuous numerical values for a target variable. The popular metrics for regression include mean absolute error/MAE, mean squared error/MSE, and R-squared/ R^2 .

Consideration of randomness is crucial when dealing with stochastic processes or incorporating random elements such as initialization procedures in model training. To ensure reproducibility, it is encouraged to fix a random seed in model codes, while to ensure robustness, it is recommended to run the model independently with different seeds. K-fold cross-validation is commonly used to tune hyperparameter, where data are divided into k equal folds, the model is trained on k-1 folds and tested on the remaining folds.

When applying machine learning techniques in autophagy research, it is crucial to be aware of potential biases that can arise and take steps to mitigate them. One significant bias source lies in the data collection process, where inherent biases may exist due to factors such as the selection of biased groups, experimental protocols, or data acquisition methods. Among these factors, the attributes, such as gender and age, can play a vital role in autophagy research. Ensuring the data collection covers diverse and comprehensive groups not only helps in mitigating bias but also promotes relevant research on the impacts of the attributes.

Another critical aspect is the potential for algorithmic bias, where ML models may inadvertently perpetuate or amplify existing biases present in the data, leading to unfair or discriminatory outcomes. Moreover, models that are overly complex or fitted too closely to the training data may capture noise or spurious correlations, leading to overfitting. Overfit models may perform well on the training data but fail to generalize to new data, leading to biased predictions. Moreover, researchers may have preconceived notions or hypotheses about the mechanisms of autophagy, which can influence the interpretation of machine learning model outputs or the selection of features and experimental designs. To mitigate these biases, it is essential to employ strategies such as ensuring diverse and representative data collection, rigorous experiments and validation, collaborating with domain experts, and conducting independent validation experiments.

4. Recent advances

In this section, we focus on addressing the second research question, where we aim to provide an extensive review of recent advances in ML-driven autophagy research. Our discussion incorporates diverse research topics related to autophagy, including gene/protein identification, biomarker identification, drug discovery, and morphological analysis, shedding light on how ML models are widely applied to diverse

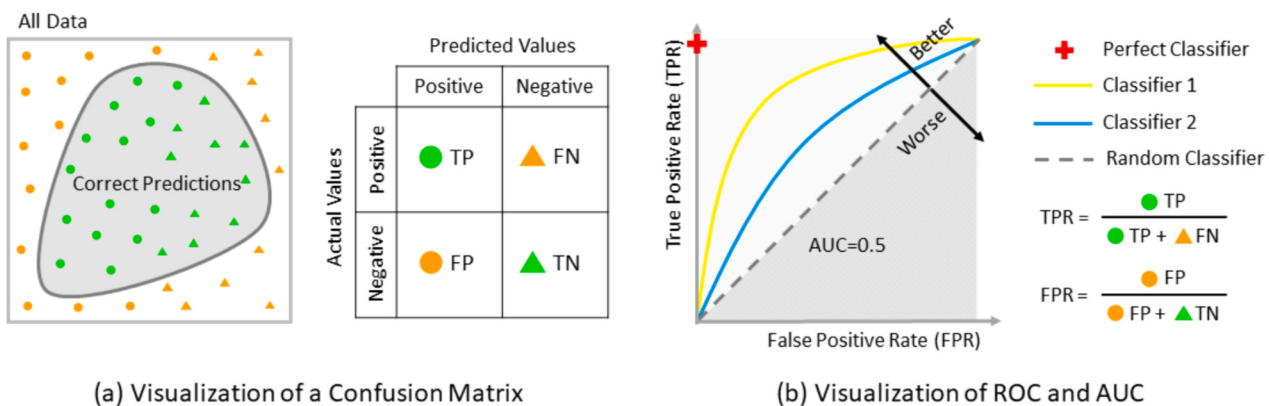


Fig. 3. Visualization of Evaluation Metrics: Confusion Matrix and ROC-AUC. (a) A visualization of a confusion matrix. The gray area indicates where the model makes correct predictions, and different shapes for predictions. We use different colors to distinguish the correct and incorrect predictions, and different shapes for positive and negative actual data. (b) Visualization of ROC and AUC. The perfect classifier has a TPR of 1, and FPR of 0. The closer the ROC curve is to the left top corner of the graph, the better the performance. Therefore, classifier 1 (yellow) is better than classifier 2 (blue) in (b).

applications. For every subsection, we begin with a brief task description, the significance, and, subsequently, the discussion on how ML has been integrated into autophagy investigations through short reviews of recent research.

4.1. Gene/protein identification

Gene/protein identification in autophagy research refers to exploring genes/proteins to identify those associated with the autophagic process. It serves as the fundamental basis for a comprehensive understanding of the molecular intricacies of autophagy, providing insights into cellular homeostasis, disease mechanisms, and potential avenues for therapeutic interventions. By expanding the gene repertoire involved in autophagy, researchers can have deeper insights into the genetic basis of autophagic regulation. Identification of autophagy-related proteins is equally critical, as understanding the functions and mechanisms of these proteins provides critical insights into the molecular machinery that governs autophagy, shedding light on how cells maintain homeostasis, respond to stress, and eliminate damaged components. These discoveries not only contribute to the broader landscape of molecular and cellular research but also hold significant implications for human health. Identification of novel autophagy-related genes or proteins can unveil potential therapeutic targets, aiding in the development of targeted interventions for diseases related to autophagic dysregulation, while also fostering innovation in the design of novel diagnostic tools and therapeutic strategies.

A general idea to tackle this task is to train ML models on genes/proteins that are known to be relevant or irrelevant to autophagy and screen new autophagy-related candidates from unknown elements. Although there has been some research conducted in this area, the topic remains underappreciated and holds significant potential. One inspiring work on gene identification [40] constructed a knowledge graph with diverse databases of genes and proteins, converted to feature vectors from the knowledge graph using the MetaPath method (i.e., to identify various network connections between proteins/genes and specific qualities), and trained an XGBoost model to predict potential autophagy-related genes. It is striking to see that 9/20 genes that were predicted in 2019 to be autophagy related have been subsequently reported as ATG genes during the intervening 3.5 years, which shows the great promise to accelerate scientific discoveries.

For protein identification, Shi et al. [41] presented ATGPred-FL to classify autophagy-related proteins and contributed a benchmark dataset of proteins. After generating informative probability features and filtering out irrelevant features, Shi et al. collected the most discriminative 52-dimensional feature to train SVM for gene classification. EnsembleDL-ATG [42] is another work that applied ensemble deep learning to combine multiple models to collaboratively predict autophagy-related proteins. Le et al. evaluated and selected the most promising model to create feature descriptors and explored the best possible combinations of base models. The final model has proven to be effective and is composed with four different deep-learning models.

4.2. Biomarker identification

In autophagy, the identification of biomarkers involves recognizing crucial genes, proteins, or other molecules associated with certain diseases. This task aids in identifying and understanding the progression of the disease and is instrumental in advancing both our understanding of cellular processes and the development of diagnostic and therapeutic strategies for the diseases. Serving as measurable indicators, biomarkers provide valuable diagnostic information and offer insights into disease progression. Identifying reliable biomarkers allows researchers and clinicians to assess the conditions of patients in various physiological and pathological conditions. Furthermore, in drug discovery, biomarkers are indispensable to assess the impact of pharmaceutical interventions on autophagic processes, facilitating the identification and optimization of

drug candidates.

Depending on recent advances in biomarker identification, the procedures can vary. To begin with, we present a general workflow of identifying autophagy-related genes as disease biomarkers in recent studies [43–54], with an overview in Fig. 4. Initially, genes associated with a specific disease or cellular process related to the disease are gathered from publications, databases, or through weighted gene co-expression network analysis/WGCNA, a method to mine gene co-expression patterns. Subsequently, differentially expressed genes (DEGs) are identified using data from healthy individuals and patients, commonly implemented with limma package in R software. The intersection between DEGs and ATG genes is considered to represent meaningful autophagy-related genes (ARGs) for the specific disease or cellular process. LASSO regularization or other feature selection approaches can be employed to rank and identify the most important features. Furthermore, ML tools for classification, such as RF, LR, SVM, or SVM-RFE, can be utilized to train a classifier with gene expression matrices to identify potential gene biomarkers. Another common subtask in biomarker identification is disease subtype discovery. Unsupervised clustering is often used to reveal potential group divisions in the data [43–45]. Such a workflow is commonly used in finding biomarkers of cancers [43,44,46–48,55], cardiovascular [45,49,50], respiratory [51,52], and other diseases [53,54,56].

Slightly different from the workflow we discussed above, Serrano et al. [57] used SVM, RF, and LR to validate if the expression of different mRNA from peripheral blood mononuclear cell can be used to classify HIV-infected patients, and recognized *GAPARAPL1* as the typical mRNA for an HIV biomarker. Another line of research [58] applied KNN on immunohistochemical images of cells to identify key ATG proteins that can be used to classify tissues with different subtypes of renal cell carcinoma (RCC), including MAP1LC3B/LC3B for clear cell RCC, or ATG5 and SQSTM1/p62 for chromophobe RCC. Park et al. [59] used the unsupervised multi-omics method, Multi-Omics Factor Analysis+/MOFA+ [60], to identify new subtypes of Alzheimer disease. In another study [61], SVM is applied to identify biological elements that maximize discrimination of different patient groups of chronic kidney disease by using transcriptomic profiles. Compared to other work focusing on accuracy, Yones et al. [62] considered interpretability by applying interpretable ML methods, such as rule-based ML models and rule networks, to offer classification transparency of diagnosis of systemic lupus erythematosus.

4.3. Drug discovery

Drug discovery refers to the identification and development of new therapeutic compounds that have the potential to alleviate, treat, or prevent various diseases. Traditional drug discovery is a complex and challenging process characterized by inherent intricacies and uncertainties, requiring substantial domain expert supervision, as well as financial and time commitment. ML plays an increasingly critical role in every step of drug discovery and is too broad of a topic for this review to cover. In this subsection, we specifically focus on the role of ML tools in the identification of new compounds.

Autophagy plays a vital role in maintaining cellular homeostasis and responding to various stressors. Dysregulation of autophagy is implicated in a range of diseases, including neurodegenerative disorders, cancer, etc. Therefore, drug discovery based on autophagy modulation has attracted attention for therapeutic interventions in various diseases. In the context of autophagy, drug discovery refers to exploring new compounds that influence the autophagy process, which can be categorized into activators and inhibitors according to their effects on enhancing or suppressing autophagy. These activators can be used to promote cell survival and clearance of toxic aggregates, whereas inhibitors block autophagy, selectively inducing cell death in certain pathological contexts. Identification of compounds that modulate autophagy provides researchers with pharmacological tools to

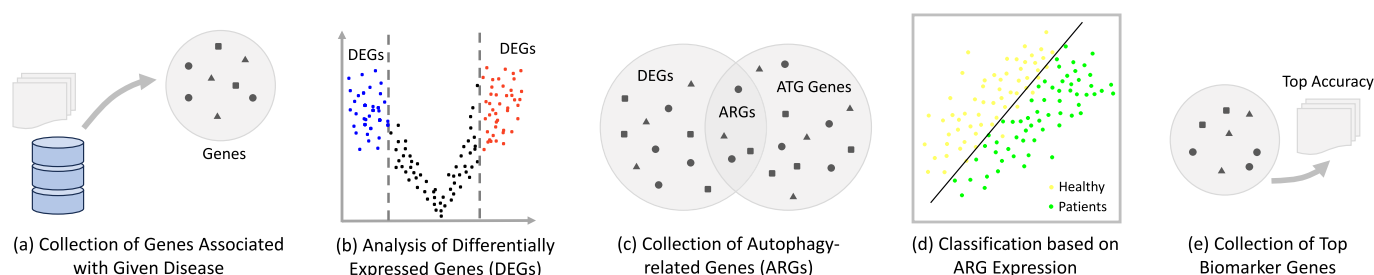


Fig. 4. General Workflow for Biomarker Gene Identification. Beginning with gene data collection, the differentially expressed genes (DEGs) are identified (red for upregulated genes, blue for downregulated genes). Autophagy-related genes (ARGs) are obtained subsequently in the intersection of DEGs and ATG genes, and their expressions are leveraged to train ML models to classify healthy individuals and patients. Top ARGs with high accuracy are collected and identified as key biomarker genes for the given diseases.

understand the autophagic process and holds promise for the treatment of diseases where autophagic dysregulation is a contributing factor. Furthermore, the identification of autophagy-modulating compounds and their interactions contribute to our understanding of the underlying molecular mechanisms of autophagy regulation. This knowledge not only aids in the development of targeted therapies but also opens avenues for the exploration of novel treatment strategies for a variety of medical conditions.

Emerging ML techniques are increasingly being integrated into drug discovery to streamline the identification of promising drug candidates. Despite applying diverse molecular representations and ML techniques, recent work [63–66] generally retrieves a collection of compounds from relevant databases, creates molecular representations, and utilizes ML tools to screen and filter potential compounds from numerous components such as are present in drug or chemical libraries for further wet bench experiments and efficient drug discovery.

Barardo et al. [63] presented one typical workflow on autophagy drug discovery. The representations were created from chemical descriptors of their structure and Gene Ontology (GO) terms annotated for proteins interacting with the compounds. An RF model was trained with the representations to predict new compounds. Besides, by ranking features of the RF model by importance, the critical GO terms and chemical structures were revealed. Valentini et al. [64] used the hybrid features of fingerprints and descriptors, and trained a quantitative structure-activity relationship/QSAR model to classify promising compounds. Xie et al. [65] created molecular representation with the Simplified Molecular Input Line Entry System/SMILES representation, along with Pharmacore (2D) and conformers (3D) fingerprints, which represent structural and spatial features in drug discovery. Three encoders are applied to embed data for 1D, 2D, and 3D representation embeddings, and the multi-representation molecule model is trained with the three decoders. The representations from the trained encoders and decoders were aggregated and clustered, and similarity scores for compounds were applied to identify top candidate compounds. Predicting new targets for autophagy modulators is another application of ML tools to autophagy drug discovery. Shi et al. [66] identified known drug-target interactions and used probabilistic matrix factorization models to predict new interactions.

4.4. Morphological analysis

Morphological analysis in autophagy research provides a visual and structural perspective on the complex dynamics of cellular components during the autophagic process. By examining morphological changes at the subcellular or cellular level, researchers gain insights into the formation and progression of the autophagic process to facilitate other studies. Understanding the morphological aspects of autophagy is crucial for accurately characterizing autophagic events, distinguishing between different stages of the process, and discriminating variations under various physiological or pathological conditions. Moreover,

morphological analysis is instrumental in validating findings from other autophagy assays and methodologies, providing a holistic understanding of cellular responses.

Current research [67–69] on morphological analysis is associated with image processing due to its dependence on images. Most of these studies focus on high-throughput systems to accelerate relevant autophagy research. However, despite its critical role in advancing our understanding, this area has not received much attention and, as such, holds considerable potential for future breakthroughs in interdisciplinary research. In 2010, Kriston-Vizi et al. [67] proposed a pipeline of image analysis with both local adaptive and global techniques for segmentation, followed by feature extraction and reduction, and SVM to identify potential inducers of autophagy. This study demonstrated the effectiveness of screening a human alveolar cancer cell line and highlighted its potential in accelerating autophagy research. Kawaoka et al. [68] provided a high-throughput system to extract morphological parameters of the autophagy-related structures to replace manual extraction, by extracting morphological parameters from fluorescence microscopy images with image processing procedures and applying RF to classify the autophagy-related structures into dot-shaped or elongated structures. Aside from high-throughput systems, Al Outa et al. [69] contributed a dataset for cell autophagy imaging, containing images of basal and activated autophagy states in *Drosophila*, including segmentation of cells, to assist future research of segmentation and classification of cells in basal or activated autophagy.

5. Open questions

We have witnessed the success of the integration of traditional ML techniques for autophagy research in recent studies, and we strongly think there is still a huge gap in fully harnessing the power of machine learning techniques, thus promising ample room for interdisciplinary collaboration to tackle the following open questions.

6. What are current challenges in ML for autophagy research?

Developing effective representations of complex biological data that capture the nuanced relationships between variables remains an ongoing challenge in ML-driven autophagy research, due to the diverse nature of data, including genetic, proteomic, and morphological data. Furthermore, with the increasing complexity and volume of biological data, selecting informative features and reducing dimensionality without losing critical information is crucial. ML techniques need to address the problem of dimensionality while ensuring that relevant features are retained for accurate modeling.

In addition, one important unsolved task in ML for autophagy is to deal with their lack of interpretability, which limits their utility in understanding the underlying biological mechanisms. Developing interpretable ML models that provide insights into the molecular intricacies of autophagy processes is essential for advancing scientific

understanding and facilitating translational applications.

Autophagy-related datasets often suffer from class imbalances, where certain classes of data are underrepresented. Addressing data imbalance and ensuring the availability of accurately labeled data for training ML models are critical challenges. Strategies for robust model training in the presence of imbalanced datasets need to be developed to ensure the reliability and generalizability of ML-driven autophagy research.

Achieving robust performance across diverse biological contexts and experimental conditions remains a challenge. Developing ML models that can generalize well to unseen data and effectively transfer knowledge between related tasks is essential for advancing autophagy research and its applications in disease diagnosis, prognosis, and drug discovery.

Addressing these challenges will require interdisciplinary collaborations between researchers in ML, biology, and medicine, as well as advancements in algorithm development, data curation, and experimental validation methodologies. Efforts to overcome these challenges will not only advance our understanding of autophagy processes but also pave the way for innovative diagnostic and therapeutic interventions in various diseases. In the next questions, we will discuss approaches that may advance autophagy research.

7. What advanced ML tools can be used in future research?

ML is a fast-evolving field and the ML tools that have been applied in autophagy research to date, are still relatively traditional and small-scale, and cannot reflect the state-of-the-art ML techniques. Therefore, it is natural to consider more advanced ML tools in future autophagy research. To get started, we introduce several potential use cases of various advanced ML tools, including graph neural networks, diffusion models, large language models, etc.

As we discussed in previous sections, GNNs are powerful neural networks designed specifically for processing data structured as graphs. Graphs can be flexible with nodes and edges representing various data or relationships, with the number of nodes/edges, and edges between nodes varying in different scenarios. The biological molecules or molecule interactions can be naturally represented as a graph structure, where GNNs can come into play in molecule identification, molecule interaction identification, or drug discovery. For example, graphs can be used to model a network of genes, with the genes as nodes, and the gene interactions as edges. GNNs are capable of capturing the complex interdependencies between nodes in a graph, which can be further used in diverse downstream tasks, such as node classification, link prediction, and graph classification. In molecule interactions for drug or biomarker discovery, the task can be formulated as link prediction, with the molecules as nodes, and the known interactions as edges, to predict new interactions between molecules. Other than edge classification, node classification can be promising in these tasks as well, learning from edges and nodes with labels, to predict new molecules and genes that are very likely to be relevant to diseases or biological processes. GNNs have shown considerable promise in the analysis of biological data. The geneDRAGNN [70] is a framework leveraging GNNs to predict gene-disease associations by integrating gene-gene interaction networks, genetic information, and additional domain knowledge including gene ontology and tissue-specific gene expression.

Generative models are able to generate new data by learning and effectively sampling from a synthetic data distribution that resembles the real data distribution and bring in the recent explosion of generative artificial intelligence tools, such as diffusion models [24–27] and LLMs [28–32]. Diffusion models trained on text and image pairs have achieved incredible performance in generating images with prompts input by users [71–73]; LLMs, such as GPT-3 [32], trained on a tremendous amount of text from the internet, can understand and answer questions on many challenging topics.

Generative models open exciting new opportunities in autophagy or

more broadly, the biomedical domain. For example, forming a hypothesis is the first step of any scientific discovery and the capability of generative models to integrate data sources in a humanly impossible scale can be harvested to generate a hypothesis. One unique advantage of hypothesis generation is that the hypothesis should undergo rigorous validation, regardless of its origin – be it human or AI. Thus, the black-box nature of generative models does not bring extra uncertainty. Another exciting line of research is to use generative models to predict how cells respond to unseen chemical or genetic perturbations, an approach whose effectiveness is already attested to by recent studies [74–78]. Autophagy malfunction is known to have serious implications in global gene expressions, but their interplay is very complicated. The impressive biological response predictive power of generative models can be potentially deployed in this direction.

Below we review diffusion models and LLMs as well as how they can be applied in a biomedical context.

Diffusion models excel in text-to-image tasks. The general idea of diffusion models is that the training data will be added with successive Gaussian noise until becoming uniform Gaussian noise, and the diffusion models learn to recover the data by reversing the process of adding noise. After the model learns this reversal procedure, we can use the diffusion models to generate new data by denoising a randomly sampled noise. Diffusion models have been proven to have potential in protein design; RFdiffusion [79] is a diffusion model of predicting protein backbones, constructed with RoseTTAFold [80], a network of predicting protein structures, as the denoising network. This model excels in various protein design tasks, including monomer and binder design, enzyme scaffolding, and symmetric oligomer design for therapeutic and metal-binding proteins. The effectiveness of RFdiffusion was validated experimentally, confirming the structures and functions of designed proteins.

LLMs are probably the most exciting innovation in the last few years, standing out for their general-purpose capabilities and super-human question answering performances. With a transformer-based architecture, LLMs are trained on large corpora to ultimately generate coherent, contextually relevant text, and perform language-related tasks such as multi-lingual translation, text summarization, and question-answering. As discussed in previous section, BioBERT [23] is a successful application of LLM on biomedical corpora. Another approach, BioGPT [81], is trained on large-scale biomedical literature for downstream tasks of biomedical natural language processing. Besides, the sequential structure of textual data might remind us of sequential biological data, such as the sequences of amino acids or nucleic acids. We may imagine that if LLMs are used to train on numerous biological sequences, they can be the next powerful tools for understanding these sequences or generating novel sequences. There have been several studies aimed at predicting protein structure with LLMs. Lin et al. [82] applied LLMs to predict atomic-level protein structures directly from primary sequences and accelerated the process of high-resolution structure prediction. This study used this method to create the ESM Metagenomic Atlas, successfully predicting structures for over 617 million metagenomic protein sequences, including more than 225 million with high confidence, thereby providing insights into the vast diversity of natural proteins. ProGen [83] leverages LLMs to generate protein sequences with predictable functions across extensive protein families, trained on 280 million sequences from over 19,000 families and enhanced with tags specifying protein properties. By predicting the next amino acid in a sequence, the optimization is performed iteratively without needing explicit structural information. Learning from a vast protein sequence database allows ProGen to learn universal representations of proteins, encompassing both local and global structure motifs.

Notably, we strongly think morphological analysis is critical but underexplored in current research, especially for the future construction of high-throughput systems to examine the cells for autophagy status. Existing efforts are still quite restrictive in the ML tool selection and lack flexibility in how to integrate ML tools in every step of the workflow.

Spatial transcriptomics that allows morphological and molecular measurement of the same cells is a very promising application domain for ML [84–86]. For example, to handle the challenge of high-throughput demand, modern techniques of image processing and computer vision can be applied to accelerate the classification and segmentation of autophagic cells or status. Another example is that existing literature proves ML effectively predicts morphology from gene expression and vice versa at scale, which has the potential to be harnessed to link morphological changes with molecular changes related to the autophagy process.

To conclude, one potential future research direction for ML in autophagy should focus on incorporating advanced ML tools such as Graph Neural Networks, diffusion models, and large language models. These tools offer enhanced capabilities for processing complex biological data, modeling molecular interactions, and generating hypotheses, thereby advancing our understanding of autophagic processes.

8. Are ML results reliable?

As many machine learning tools are becoming increasingly complex in their design and mysterious in their internal mechanism, they are often dubbed as black-box models. Interpreting models, especially deep-learning models, poses several challenges due to their complex architectures and high-dimensional feature representations. Additionally, deep learning models often involve millions of parameters, making it challenging to identify which features are most influential for a particular prediction. There is a growing interest in understanding the decision-making process and improving the model interpretability and transparency. In some low-risk scenarios of ML applications, such as predicting if a new gene is related to autophagy, we might have many candidate genes to screen and want to start with genes that are most likely to be the target genes. In such cases, we might be more interested in the likelihood predicted by a model with guaranteed reliable performance for savings in resources. In other scenarios where reliability is of great importance, such as a situation where a false prediction would bring significant repercussions, we also want to learn why and when the model tends to make false decisions to help us use it cautiously. In both scenarios, models with better explainability and transparency are preferable.

Interpretability and explainability are two terms that are commonly used interchangeably when assessing model transparency. However, they do have intricate differences.

Interpretability often refers to intrinsic transparency for models, which means we can understand the underlying mechanism of the decision process. In practice, there is usually a trade-off between model complexity and interpretability as models with high interpretability often have limited flexibility. Typically, we consider linear models and decision trees as the most interpretable models. Linear models can be interpreted by the weights of the model but are unable to model or even self-learn nonlinear relationships. Decision trees are comprised of a sequence of binary rules to segment data that are easy for humans to understand. In contrast, the ensemble of decision trees, i.e. RF, is always considered as non-interpretable, because the prediction is voted among multiple trees thus making it difficult to interpret the decisions.

Explainability focuses on explaining the behaviors of complex models by probing with other standalone models. Instead of seeking to fully understand the inner mechanics of the models, we discover the relationship between input data and model outputs with model-agnostic methods. For example, identifying feature importance [43,45,46,50–52,55,56] is critical to explain the model and has been commonly used in current ML-driven autophagy research. RF, LASSO Cox regression, and XGBoost are common methods that provide feature importance scores. The dependence of the prediction on different features can serve as an indicator of the feature importance. To this end, techniques including Local Interpretable Model-agnostic Explanations (LIME) [87] and SHapley Additive exPlanations (SHAP) [88] can be

applied to estimate the feature influences. LIME elucidates the predictions of a model by constructing a local and interpretable surrogate model and is trained on the data surrounding the specific instance of interest. Thus, LIME provides insight into decision-making in the vicinity of the instance. SHAP, in contrast, explains individual predictions based on the contribution of each feature to the prediction. SHAP calculates the Shapley value, which is a concept originating from game theory, providing a measure of feature importance and considering the interaction between features. Both methods are crucial in making complex models more transparent and understandable. The major difference between SHAP and LIME lies in their scope of explanations. SHAP provides both global explanations, capturing the overall feature importance across the entire dataset, and local explanations for individual predictions, whereas LIME focuses solely on generating local explanations by approximating the model's behavior in the vicinity of a specific instance. Compared to SHAP, LIME is more computationally efficient and thus more flexible. In practice, it is recommended to use diverse interpretation methods to gain a more comprehensive understanding of models and get domain experts involved in the interpretation results.

Similar to the interpretation of feature importance, saliency maps [89] are a technique that can highlight the regions or pixels that are most relevant to the predictions of models and provide insights for interpreting the models that use imaging data (e.g., microscopy images of cells). Moreover, another technique called counterfactual explanations [90], can be used to identify how the model's predictions would change if certain features or conditions were altered, thus understanding the causal relationships between different factors, as well as identifying potential intervention points.

Overall, interpretability has stricter criteria than explainability. Among papers we reviewed in previous sections, there are several models that are able to infer with transparency. For example, the RF and LASSO Cox regression have been applied to rank the importance of features in the prediction [43,45,46,50–52,55,56], and the rule-based models are applied for diagnosis predictions with transparency [62]. However, interpretability and explainability are still largely overlooked due to model transparency-utility tradeoff in interdisciplinary research. Biomedicine is a domain where errors such as with true positive can have significant implications, thus rarely relying on black-box models. We call for more emphasis on this aspect in order for more advanced models to be recognized and received by the broader biomedical community.

Apart from enhancing interpretable and explainable ML tools, it would be more reliable if the predictions could be validated from rigorous scientific experiments. Of note, ML outcomes remain predictive in nature, and while recent advancements demonstrate promising outcomes, experimental validation remains crucial and should be implemented in tandem to reinforce findings. In practice, it is recommended to use ML techniques as auxiliary tools instead of main contributors.

9. Conclusion

As an essential process in cellular maintenance, autophagy underscores its significance with its complicated involvement in numerous biological processes and human diseases. The emergence of studies applying ML tools to assist in revealing the internal mechanism of autophagy showcases the promise of interdisciplinary collaborations. This article does not intend to exhaustively review all literature involving the use of ML techniques to investigate the autophagy process. Instead, its goal is to provide a taxonomy of ML models that are applicable to this field of study, with a particular emphasis on highlighting recent progress in generative modeling and its greater potential in a biomedical context. We advocate for a reciprocal curiosity between ML and autophagy, to comprehend challenges each field presents. Our aim is to break the interdisciplinary barrier and find solutions to many open questions through collaborations.

CRediT authorship contribution statement

Ying Yang: Writing – review & editing, Writing – original draft, Conceptualization. **Zhaoying Pan:** Writing – original draft, Conceptualization. **Jianhui Sun:** Writing – original draft, Conceptualization. **Joshua Welch:** Writing – original draft, Supervision, Conceptualization. **Daniel J. Klionsky:** Writing – review & editing, Supervision.

Declaration of competing interest

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

Data availability

No data was used for the research described in the article.

Acknowledgements

This work was supported by NIH grant GM131919 to DJK.

References

- [1] D.J. Klionsky, The molecular machinery of autophagy: unanswered questions, *J. Cell Sci.* 118 (Pt 1) (2005) 7.
- [2] Y. Chen, D.J. Klionsky, The regulation of autophagy—unanswered questions, *J. Cell Sci.* 124 (2) (2011) 161–170.
- [3] Y. Yang, D.J. Klionsky, Autophagy and disease: unanswered questions, *Cell Death Differ.* 27 (3) (2020) 858–871.
- [4] D. Glick, S. Barth, K.F. Macleod, Autophagy: cellular and molecular mechanisms, *J. Pathol.* 221 (1) (2010) 3–12.
- [5] K.R. Parzych, D.J. Klionsky, An overview of autophagy: morphology, mechanism, and regulation, *Antioxid. Redox Signal.* 20 (3) (2014) 460–473.
- [6] Y. Feng, D. He, Z. Yao, D.J. Klionsky, The machinery of macroautophagy, *Cell Res.* 24 (1) (2014) 24–41.
- [7] D. Gatica, V. Lahiri, D.J. Klionsky, Cargo recognition and degradation by selective autophagy, *Nat. Cell Biol.* 20 (3) (2018) 233–242.
- [8] L. Wang, D.J. Klionsky, H.-M. Shen, The emerging mechanisms and functions of microautophagy, *Nat. Rev. Mol. Cell Biol.* 24 (3) (2023) 186–203.
- [9] B. Loos, D.J. Klionsky, E. Wong, Augmenting brain metabolism to increase macro- and chaperone-mediated autophagy for decreasing neuronal proteotoxicity and aging, *Prog. Neurobiol.* 156 (2017) 90–106.
- [10] I.H. Witten, E. Frank, M.A. Hall, C.J. Pal, M. Data, Practical machine learning tools and techniques, in: *Data Mining vol. 2*, no. 4, Elsevier Amsterdam, The Netherlands, 2005, pp. 403–413.
- [11] R. Tibshirani, The lasso method for variable selection in the Cox model, *Stat. Med.* 16 (4) (1997) 385–395.
- [12] E. Fix, J.L. Hodges, Discriminatory analysis. Nonparametric discrimination: consistency properties, *International Statistical Review/Revue Internationale de Statistique* 57 (3) (1989) 238–247.
- [13] T. Cover, P. Hart, Nearest neighbor pattern classification, *IEEE Trans. Inf. Theory* 13 (1) (1967) 21–27.
- [14] T.K. Ho, Random decision forests, in: *Proceedings of 3rd International Conference on Document Analysis and Recognition vol. 1*, IEEE, 1995, pp. 278–282.
- [15] L. Breiman, Random forests, *Mach. Learn.* 45 (2001) 5–32.
- [16] T. Chen, C. Guestrin, Xgboost: a scalable tree boosting system, in: *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 2016, pp. 785–794.
- [17] C. Cortes, V. Vapnik, Support-vector networks, *Mach. Learn.* 20 (1995) 273–297.
- [18] I. Guyon, J. Weston, S. Barnhill, V. Vapnik, Gene selection for cancer classification using support vector machines, *Mach. Learn.* 46 (2002) 389–422.
- [19] Y. LeCun, et al., Backpropagation applied to handwritten zip code recognition, *Neural Comput.* 1 (4) (1989) 541–551.
- [20] D.E. Rumelhart, G.E. Hinton, R.J. Williams, Learning Internal Representations by Error Propagation, Institute for Cognitive Science, University of California, San Diego La, 1985.
- [21] F. Scarselli, M. Gori, A.C. Tsoi, M. Hagenbuchner, G. Monfardini, The graph neural network model, *IEEE Trans. Neural Netw.* 20 (1) (2008) 61–80.
- [22] S. Jaeger, S. Fulle, S. Turk, Mol2vec: unsupervised machine learning approach with chemical intuition, *J. Chem. Inf. Model.* 58 (1) (2018) 27–35.
- [23] J. Lee, et al., BioBERT: a pre-trained biomedical language representation model for biomedical text mining, *Bioinformatics* 36 (4) (2020) 1234–1240.
- [24] J. Sohl-Dickstein, E. Weiss, N. Maheswaranathan, S. Ganguli, Deep unsupervised learning using nonequilibrium thermodynamics, in: *International Conference on Machine Learning*, PMLR, 2015, pp. 2256–2265.
- [25] Y. Song, J. Sohl-Dickstein, D.P. Kingma, A. Kumar, S. Ermon, B. Poole, Score-based generative Modeling Through Stochastic Differential Equations, *arXiv preprint arXiv:2011.13456*, 2020.
- [26] J. Ho, A. Jain, P. Abbeel, Denoising diffusion probabilistic models, in: *Advances in Neural Information Processing Systems* 33, 2020, pp. 6840–6851.
- [27] J. Song, C. Meng, S. Ermon, Denoising Diffusion Implicit Models, *arXiv preprint arXiv:2010.02502*, 2020.
- [28] A. Vaswani, et al., Attention is all you need, in: *Advances in Neural Information Processing Systems* 30, 2017.
- [29] J. Devlin, M.-W. Chang, K. Lee, K. Toutanova, Bert: Pre-training of Deep Bidirectional Transformers for Language Understanding, *arXiv preprint arXiv:1810.04805*, 2018.
- [30] A. Radford, K. Narasimhan, T. Salimans, I. Sutskever, Improving Language Understanding by Generative Pre-training, 2018.
- [31] A. Radford, J. Wu, R. Child, D. Luan, D. Amodei, I. Sutskever, Language models are unsupervised multitask learners, *OpenAI blog* 1 (8) (2019) 9.
- [32] T. Brown, et al., Language models are few-shot learners, in: *Advances in Neural Information Processing Systems* 33, 2020, pp. 1877–1901.
- [33] O. Ronneberger, P. Fischer, T. Brox, U-net: convolutional networks for biomedical image segmentation, in: *Medical Image Computing and Computer-Assisted Intervention—MICCAI 2015: 18th International Conference*, Munich, Germany, October 5–9, 2015, Proceedings, Part III 18, Springer, 2015, pp. 234–241.
- [34] K. He, X. Zhang, S. Ren, J. Sun, Deep residual learning for image recognition, in: *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2016, pp. 770–778.
- [35] S. Hochreiter, J. Schmidhuber, Long short-term memory, *Neural Comput.* 9 (8) (1997) 1735–1780.
- [36] K. Cho, et al., Learning Phrase Representations Using RNN Encoder-Decoder for Statistical Machine Translation, *arXiv preprint arXiv:1406.1078*, 2014.
- [37] T.N. Kipf, M. Welling, Semi-supervised Classification with Graph Convolutional Networks, *arXiv preprint arXiv:1609.02907*, 2016.
- [38] P. Veličković, G. Cucurull, A. Casanova, A. Romero, P. Lio, Y. Bengio, Graph Attention Networks, *arXiv preprint arXiv:1710.10903*, 2017.
- [39] T. Mikolov, K. Chen, G. Corrado, J. Dean, Efficient Estimation of Word Representations in Vector Space, *arXiv preprint arXiv:1301.3781*, 2013.
- [40] T.I. Oprea, J.J. Yang, D.R. Byrd, V. Deretic, Autophagy Dark Genes: Can We Find Them with Machine Learning?, *bioRxiv*, 2019 715037.
- [41] S. Jiao, Z. Chen, L. Zhang, X. Zhou, L. Shi, ATGPred-FL: sequence-based prediction of autophagy proteins with feature representation learning, *Amino Acids* 54 (5) (2022) 799–809.
- [42] L. Yu, Y. Zhang, L. Xue, F. Liu, R. Jing, J. Luo, EnsembleDL-ATG: identifying autophagy proteins by integrating their sequence and evolutionary information using an ensemble deep learning framework, *Comput. Struct. Biotechnol. J.* 21 (2023) 4836–4848.
- [43] Q. Wei, X. Jiang, X. Miao, Y. Zhang, F. Chen, P. Zhang, Molecular subtypes of lung adenocarcinoma patients for prognosis and therapeutic response prediction with machine learning on 13 programmed cell death patterns, *J. Cancer Res. Clin. Oncol.* 149 (13) (2023) 11351–11368.
- [44] C. Wei, et al., Machine learning and single-cell sequencing reveal the potential regulatory factors of mitochondrial autophagy in the progression of gastric cancer, *J. Cancer Res. Clin. Oncol.* 149 (17) (2023) 15561–15572.
- [45] J. Lv, D. Wang, T. Li, Autophagy-mediated expression clusters are involved in immunity regulation of coronary artery disease, *BMC Genomic Data* 23 (1) (2022) 1–14.
- [46] J. Ding, C. Wang, Y. Sun, J. Guo, S. Liu, Z. Cheng, Identification of an autophagy-related signature for prognosis and immunotherapy response prediction in ovarian cancer, *Biomolecules* 13 (2) (2023) 339.
- [47] G.-Z. Zhang, et al., Development of a machine learning-based autophagy-related lncrna signature to improve prognosis prediction in osteosarcoma patients, *Front. Mol. Biosci.* 8 (2021) 615084.
- [48] W. Liu, et al., FOLFOX treatment response prediction in metastatic or recurrent colorectal cancer patients via machine learning algorithms, *Cancer Med.* 9 (4) (2020) 1419–1429.
- [49] F. Zhang, et al., Machine learning and bioinformatics to identify 8 autophagy-related biomarkers and construct gene regulatory networks in dilated cardiomyopathy, *Sci. Rep.* 12 (1) (2022) 15030.
- [50] W.-X. Yang, F.-F. Wang, Y.-Y. Pan, J.-Q. Xie, M.-H. Lu, C.-G. You, Comparison of ischemic stroke diagnosis models based on machine learning, *Front. Neurol.* 13 (2022) 1014346.
- [51] S. Xiao, et al., Identifying autophagy-related genes as potential targets for immunotherapy in tuberculosis, *Int. Immunopharmacol.* 118 (2023) 109956.
- [52] F. Xia, et al., Development of genomic phenotype and immunophenotype of acute respiratory distress syndrome using autophagy and metabolism-related genes, *Front. Immunol.* 14 (2023).
- [53] G. Dong, H. Gao, Y. Chen, H. Yang, Machine learning and bioinformatics analysis to identify autophagy-related biomarkers in peripheral blood for rheumatoid arthritis, *Front. Genet.* 14 (2023).
- [54] Y. Deng, et al., Machine learning models identify ferroptosis-related genes as potential diagnostic biomarkers for Alzheimer's disease, *Front. Aging Neurosci.* 14 (2022) 994130.
- [55] R.F. Zaarour, et al., Genomic analysis of waterpipe smoke-induced lung tumor autophagy and plasticity, *Int. J. Mol. Sci.* 23 (12) (2022) 6848.
- [56] X. Wang, Z. Wang, Z. Guo, Z. Wang, F. Chen, Z. Wang, Exploring the role of different cell-death-related genes in sepsis diagnosis using a machine learning algorithm, *Int. J. Mol. Sci.* 24 (19) (2023) 14720.
- [57] A. Serrano, et al., Dysregulation of apoptosis and autophagy gene expression in peripheral blood mononuclear cells of efficiently treated HIV-infected patients, *Aids* 32 (12) (2018) 1579–1587.

- [58] Z. He, H. Liu, H. Moch, H.-U. Simon, Machine learning with autophagy-related proteins for discriminating renal cell carcinoma subtypes, *Sci. Rep.* 10 (1) (2020) 720.
- [59] J.C. Park, et al., Multi-omics-based autophagy-related untypical subtypes in patients with cerebral amyloid pathology, *Adv. Sci.* 9 (23) (2022) 2201212.
- [60] R. Argelaguet, et al., Multi-omics factor analysis—a framework for unsupervised integration of multi-omics data sets, *Mol. Syst. Biol.* 14 (6) (2018) e8124.
- [61] S. Granata, et al., Autophagy Activation in Peripheral Blood Mononuclear Cells of Peritoneal Dialysis Patients, *Kidney International Reports*, 2023.
- [62] S.A. Yones, et al., Interpretable machine learning identifies paediatric systemic lupus erythematosus subtypes based on gene expression data, *Sci. Rep.* 12 (1) (2022) 7433.
- [63] D.G. Barardo, D. Newby, D. Thornton, T. Ghafourian, J.P. de Magalhães, A. Freitas, Machine learning for predicting lifespan-extending chemical compounds, *Aging (Albany NY)* 9 (7) (2017) 1721.
- [64] E. Valentini, et al., Targeting the anti-apoptotic Bcl-2 family proteins: machine learning virtual screening and biological evaluation of new small molecules, *Theranostics* 12 (5) (2022) 2427.
- [65] C. Xie, et al., Amelioration of Alzheimer's disease pathology by mitophagy inducers identified via machine learning and a cross-species workflow, *Nat. Biomed. Eng.* 6 (1) (2022) 76–93.
- [66] Q. Shi, et al., Mechanisms of action of autophagy modulators dissected by quantitative systems pharmacology analysis, *Int. J. Mol. Sci.* 21 (8) (2020) 2855.
- [67] J. Kriston-Vizi, C.A. Lim, P. Condron, K. Chua, M. Wasser, H. Flotow, An automated high-content screening image analysis pipeline for the identification of selective autophagic inducers in human cancer cell lines, *J. Biomol. Screen.* 15 (7) (2010) 869–881.
- [68] T. Kawaoka, S. Ohnuki, Y. Ohya, K. Suzuki, Morphometric analysis of autophagy-related structures in *Saccharomyces cerevisiae*, *Autophagy* 13 (12) (2017) 2104–2110.
- [69] A. Al Outa, et al., Cellular, a cell autophagy imaging dataset, *Scientific Data* 10 (1) (2023) 806.
- [70] A. Altaab, D. Huang, C. Byles-Ho, H. Khatib, F. Sosa, T. Hu, geneDRAGNN: gene disease prioritization using graph neural networks, in: 2022 IEEE Conference on Computational Intelligence in Bioinformatics and Computational Biology (CIBCB), IEEE, 2022, pp. 1–10.
- [71] A. Nichol, et al., Glide: Towards Photorealistic Image Generation and Editing with Text-guided Diffusion Models, *arXiv preprint arXiv:2112.10741*, 2021.
- [72] R. Rombach, A. Blattmann, D. Lorenz, P. Esser, B. Ommer, High-resolution image synthesis with latent diffusion models, in: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition, 2022, pp. 10684–10695.
- [73] C. Saharia, et al., Photorealistic text-to-image diffusion models with deep language understanding, *Adv. Neural Inf. Proces. Syst.* 35 (2022) 36479–36494.
- [74] M. Lotfollahi, et al., Predicting cellular responses to complex perturbations in high-throughput screens, *Mol. Syst. Biol.* 19 (2023) e11517.
- [75] J. Welch, Multiplying insights from perturbation experiments: predicting new perturbation combinations, *Mol. Syst. Biol.* 19 (2023) e11667.
- [76] Y. Roohani, K. Huang, J. Leskovec, Predicting transcriptional outcomes of novel multigene perturbations with gears, *Nat. Biotechnol.* (2023) 1–9.
- [77] Z. Piran, N. Cohen, Y. Hoshen, M. Nitzan, Disentanglement of single-cell data with biolord, *Nat. Biotechnol.* (2024) 1–6.
- [78] H. Yu, J.D. Welch, Perturbnet Predicts Single-cell Responses to Unseen Chemical and Genetic Perturbations, *BioRxiv*, p. 2022.07. 20.500854, 2022.
- [79] J.L. Watson, et al., De novo design of protein structure and function with RFdiffusion, *Nature* 620 (7976) (2023) 1089–1100.
- [80] M. Baek, et al., Accurate prediction of protein structures and interactions using a three-track neural network, *Science* 373 (6557) (2021) 871–876.
- [81] R. Luo, et al., BioGPT: generative pre-trained transformer for biomedical text generation and mining, *Brief. Bioinform.* 23 (6) (2022) bbac409.
- [82] Z. Lin, et al., Evolutionary-scale prediction of atomic-level protein structure with a language model, *Science* 379 (6637) (2023) 1123–1130.
- [83] A. Madani, et al., Large language models generate functional protein sequences across diverse families, *Nat. Biotechnol.* (2023) 1–8.
- [84] K.J. Kobayashi-Kirschvink, et al., Prediction of single-cell RNA expression profiles in live cells by Raman microscopy with Raman2RNA, *Nat. Biotechnol.* (2024) 1–9.
- [85] D. Zhang, et al., Inferring super-resolution tissue architecture by integrating spatial transcriptomics with histology, *Nat. Biotechnol.* (2024) 1–6.
- [86] H. Lee, J.D. Welch, MorphNet Predicts Cell Morphology from Single-cell Gene Expression, *bioRxiv*, p. 2022.10. 21.513201, 2022.
- [87] M.T. Ribeiro, S. Singh, C. Guestrin, “Why should i trust you?” Explaining the predictions of any classifier, in: Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, 2016, pp. 1135–1144.
- [88] S.M. Lundberg, S.-I. Lee, A unified approach to interpreting model predictions, *Adv. Neural Inf. Proces. Syst.* 30 (2017).
- [89] K. Simonyan, A. Vedaldi, A. Zisserman, Deep Inside Convolutional Networks: Visualising Image Classification Models and Saliency Maps, *arXiv preprint arXiv:1312.6034*, 2013.
- [90] S. Wachter, B. Mittelstadt, C. Russell, Counterfactual explanations without opening the black box: automated decisions and the GDPR, *Harv. J.L. & Tech.* 31 (2017) 841.