

Abstract

Modern machine learning systems are designed to learn from vast amounts of data, necessitating massively over-parameterized models. As these systems grow in scale, they become prohibitively expensive to train and deploy while simultaneously accumulating sensitive information about their training data. This retention of data-dependent knowledge exposes models to privacy breaches, adversarial manipulation, and legal challenges such as the right to be forgotten. Consequently, the ability to make models not only learn efficiently but also forget selectively and securely has become a fundamental requirement in modern AI. This thesis presents a unified framework for efficient learning and forgetting, exploring how models can be trained faster with fewer resources and later untrained safely without compromising accuracy or robustness.

Since machine learning arises from the interaction between two key components, data (inputs) and models (parameters), reducing input size and parameter count can yield substantial gains in both training and inference time while also mitigating the impact of noise. The first contribution explores this idea in the feature space through Predictive Permutation Feature Selection (PPFS), a novel Markov blanket-based wrapper feature selection algorithm that identifies the smallest subset of features retaining full predictive power. PPFS introduces a new conditional independence test, Predictive Permutation Independence (PPI), enabling accurate and efficient feature selection for both categorical and continuous variables across regression and classification tasks.

Recognizing that the input bottleneck extends beyond features to the volume of data itself, the second contribution focuses on data pruning with a Robust, Class-Aware Probabilistic (RCAP) dynamic pruning algorithm that adaptively samples informative training examples each epoch, prioritizing minority-class performance.

RCAP achieves over $8\times$ speedups with negligible accuracy loss, demonstrating that efficient training need not compromise robustness.

Having addressed input redundancy, the thesis then turns to model pruning through Structured Pruning via Ranking (SPvR), which eliminates redundant neurons, filters, and even layers without backpropagation using a novel ranking and grouping function. SPvR simultaneously reduces model depth and width, producing compact sub-networks that train faster and infer significantly quicker, making them ideal for edge and resource-constrained environments.

Finally, once efficient models are deployed, the question of secure forgetting becomes essential. To this end, a Scalable, Adaptive, and Label-Agnostic machine unlearning algorithm (SALSA) enables models to forget specific data without re-training. SALSA achieves up to $120\times$ faster unlearning than existing methods while remaining resilient to privacy and adversarial attacks.

Together, these contributions trace a cohesive journey from efficient learning to responsible forgetting, laying the foundation for scalable, adaptive, and privacy-conscious machine intelligence.

Keywords: Feature Selection, Data Pruning, Structured Pruning, Machine Unlearning.