
A Benchmark for Interpretability Methods in Deep Neural Networks

Sara Hooker, Dumitru Erhan, Pieter-Jan Kindermans, Been Kim
Google Brain
shooker,dumitru,pikinder,beenkim@google.com

Abstract

We propose an empirical measure of the approximate accuracy of feature importance estimates in deep neural networks. Our results across several large-scale image classification datasets show that many popular interpretability methods produce estimates of feature importance that are not better than a random designation of feature importance. Only certain ensemble based approaches—VarGrad and SmoothGrad-Squared—outperform such a random assignment of importance. The manner of ensembling remains critical, we show that some approaches do no better than the underlying method but carry a far higher computational burden.

1 Introduction

In a machine learning setting, a question of great interest is estimating the influence of a given input feature to the prediction made by a model. Understanding what features are important helps improve our models, builds trust in the model prediction and isolates undesirable behavior. Unfortunately, it is challenging to evaluate whether an explanation of model behavior is reliable. First, there is no ground truth. If we knew what was important to the model, we would not need to estimate feature importance in the first place. Second, it is unclear which of the numerous proposed interpretability methods that estimate feature importance one should select [7, 6, 44, 31, 38, 34, 40, 37, 20, 23, 12, 10, 41, 32, 42, 28, 35, 2]. Many feature importance estimators have interesting theoretical properties e.g. preservation of relevance [6] or implementation invariance [38]. However even these methods need to be configured correctly [23, 38] and it has been shown that using the wrong configuration can easily render them ineffective [19]. For this reason, it is important that we build a framework to empirically validate the relative merits and reliability of these methods.

A commonly used strategy is to remove the supposedly informative features from the input and look at how the classifier degrades [30]. This method is cheap to evaluate but comes at a significant drawback. Samples where a subset of the features are removed come from a different distribution (as can be seen in Fig. 1). Therefore, this approach clearly violates one of the key assumptions in machine learning: the training and evaluation data come from the same distribution. Without re-training, it is unclear whether the degradation in model performance comes from the distribution shift or because the features that were removed are truly informative [10, 12].

For this reason we decided to verify how much information can be removed in a typical dataset before accuracy of a retrained model breaks down completely. In this experiment, we applied ResNet-50 [17], one of the most commonly used models, to ImageNet. It turns out that removing information is quite hard. With 90% of the inputs removed the network still achieves 63.53% accuracy compared to 76.68% on clean data. This implies that a strong performance degradation without re-training *might* be caused by a shift in distribution instead of removal of information.

Instead, in this work we evaluate interpretability methods by verifying how the accuracy of a retrained model degrades as features estimated to be important are removed. We term this approach **ROAR**, **ReMOve And Retrain**. For each feature importance estimator, ROAR replaces the fraction of the

pixels estimated to be most important with a fixed uninformative value. This modification (shown in Fig. 1) is repeated for each image in both the training and test set. To measure the change to model behavior after the removal of these input features, we separately train new models on the modified dataset such that train and test data comes from a similar distribution. More accurate estimators will identify as important input pixels whose subsequent removal causes the sharpest degradation in accuracy. We also compare each method performance to a *random* assignment of importance and a sobel edge filter [36]. Both of these control variants produce rankings that are independent of the properties of the model we aim to interpret. Given that these methods do not depend upon the model, the performance of these variants represent a lower bound of accuracy that a interpretability method could be expected to achieve. In particular, a random baseline allows us to answer the question: *is the interpretability method more accurate than a random guess as to which features are important?* In Section 3 we will elaborate on the motivation and the limitations of ROAR.

We applied ROAR in a broad set of experiments across three large scale, open source image datasets: ImageNet [11], Food 101 [9] and Birdsnap [8]. In our experiments we show the following.

- Training performance is quite robust to removing input features. For example, after randomly replacing 90% of all ImageNet input features, we can still train a model that achieves 63.53 ± 0.13 (average across 5 independent runs). This implies that a small subset of features are sufficient for the actual decision making. Our observation is consistent across datasets.
- The base methods we evaluate are no better or on par with a random estimate at finding the core set of informative features. However, we show that SmoothGrad-Squared (an unpublished variant of Classic SmoothGrad [35]) and Vargrad [2], methods which ensemble a set of estimates produced by basic methods, far outperform both the underlying method and a random guess. These results are consistent across datasets and methods.
- Not all ensemble estimators improve performance. Classic SmoothGrad [35] is worse than a single estimate despite being more computationally intensive.

2 Related Work

Interpretability research is diverse, and many different approaches are used to gain intuition about the function implemented by a neural network. For example, one can distill or constrain a model into a functional form that is considered more interpretable [5, 13, 39, 29]. Other methods explore the role of neurons or activations in hidden layers of the network [25, 27, 24, 43], while others use high level concepts to explain prediction results [18]. Finally there are also the input feature importance estimators that we evaluate in this work. These methods estimate the importance of an input feature for a specified output activation.

While there is no clear way to measure “correctness”, comparing the relative merit of different estimators is often based upon human studies [31, 28, 22] which interrogate whether the ranking is meaningful to a human. Recently, there have been efforts to evaluate whether interpretability methods are *both* reliable and meaningful to human. For example, in [19] a unit test for interpretability methods is proposed which detects whether the explanation can be manipulated by factors that are not affecting the decision making process. Another approach considers a set of sanity checks that measure the change to an estimate as parameters in a model or dataset labels are randomized [3]. Closely related to this manuscript are the modification-based evaluation measures proposed originally by [30] with subsequent variations [20, 26]. In this line of work, one replaces the inputs estimated to be most important with a value considered meaningless to the task. These methods measure the subsequent degradation to the trained model at inference time. Recursive feature elimination methods [16] are a greedy search where the algorithm is trained on an iteratively altered subset of features. Recursive feature elimination does not scale to high dimensional datasets (one would have to retrain removing one pixel at a time) and unlike our work is a method to estimate feature importance (rather than evaluate different existing interpretability methods).

To the best of our knowledge, unlike prior modification based evaluation measures, our benchmark requires retraining the model from random initialization on the modified dataset rather than re-scoring the modified image at inference time. Without this step, we argue that one cannot decouple whether the model’s degradation in performance is due to artifacts introduced by the value used to replace

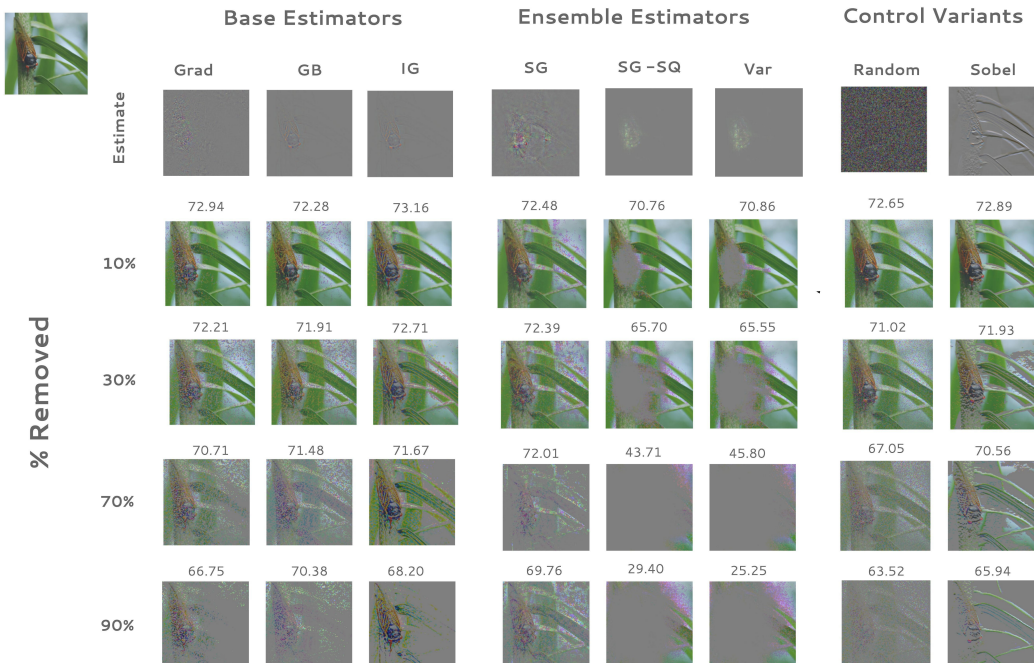


Figure 1: A single ImageNet image modified according to the ROAR framework. The fraction of pixels estimated to be most important by each interpretability method is replaced with the mean. Above each image, we include the average test-set accuracy for 5 ResNet-50 models independently trained on the modified dataset. **From left to right:** base estimators (gradient heatmap (GRAD), Integrated Gradients (IG), Guided Backprop (GB)), derivative approaches that ensemble a set of estimates (SmoothGrad Integrated Gradients (SG-SQ-IG), SmoothGrad-Squared Integrated Gradients (SG-SQ-IG), VarGrad Integrated Gradients (Var-IG)) and control variants (random modification (Random) and a sobel edge filter (Sobel)). This image is best visualized in digital format.

the pixels that are removed or due to the approximate accuracy of the estimator. Our work considers several large scale datasets, whereas all previous evaluations have involved a far smaller subset of data [4, 30].

3 ROAR: Remove And Retrain

To evaluate a feature importance estimate using ROAR, we sort the input dimensions according to the estimated importance. We compute an estimate e of feature importance for every input in the training and test set. We rank each e into an ordered set $\{e_i^o\}_{i=1}^N$. For the top t fraction of this ordered set, we replace the corresponding pixels in the raw image with the per channel mean. We generate new train and test datasets at different degradation levels $t = [0., 10., \dots, 100]$ (where t is a percentage of all features modified). Afterwards the model is re-trained from random initialization on the new dataset and evaluated on the new test data.

Of course, because re-training can result in slightly different models, it is essential to repeat the training process multiple times to ensure that the variance in accuracy is low. To control for this, we repeat training 5 times for each interpretability method e and level of degradation t . We introduce the methodology and motivation for ROAR in the context of linear models and deep neural networks. However, we note that the properties of ROAR differ given an algorithm that explicitly uses feature selection (e.g. L1 regularization or any mechanism which limits the features available to the model at inference time). In this case one should of course mask the inputs that are known to be ignored by the model, *before* re-training. This will prevent them from being used after re-training, which could otherwise corrupt the ROAR metric. For the remainder of this paper, we focus on the performance of ROAR given deep neural networks and linear models which do not present this limitation

What would happen without re-training? The re-training is the most computationally expensive aspect of ROAR. One should question whether it is actually needed. We argue that re-training is needed because machine learning models typically assume that the train and the test data comes from a similar distribution.

The replacement value c can only be considered uninformative if the model is trained to learn it as such. Without retraining, it is unclear whether degradation in performance is due to the introduction of artifacts outside of the original training distribution or because we actually removed information. This is made explicit in our experiment in Section 4.3.1, we show that without retraining the degradation is far higher than the modest decrease in performance observed with re-training. This suggests retraining has better controlled for artefacts introduced by the modification.

Are we evaluating the right aspects? Re-training does have limitations. For one, while the architecture is the same, the model used during evaluation is not the same as the model on which the feature importance estimates were originally obtained. To understand why ROAR is still meaningful we have to think about what happens when the accuracy degrades, especially when we compare it to a random baseline. The possibilities are:

1. **We remove input dimensions and the accuracy drops.** In this case, it is very likely that the removed inputs were informative to the original model. ROAR thus gives a good indication that the importance estimate is of high quality.
2. **We remove inputs and the accuracy does not drop.** This can be explained as either:
 - (a) It could be caused by removal of an input that was uninformative to the model. This includes the case where the input might have been informative but not in a way that is useful to the model, for example, when a linear model is used and the relation between the feature and the output is non-linear. Since in such a case the information was not used by the model and it does not show in ROAR we can assume ROAR behaves as intended.
 - (b) There might be redundancy in the inputs. The same information could be represented in another feature. This behavior can be detected with ROAR as we will show in our toy data experiment.

Validating the behavior of ROAR on artificial data. To demonstrate the difference between ROAR and an approach without re-training in a controlled environment we generate a 16 dimensional dataset with 4 informative features. Each datapoint $\mathbf{x}$ and its label y was generated as follows:

$$\mathbf{x} = \frac{\mathbf{a}z}{10} + \mathbf{d}\eta + \frac{\epsilon}{10}, \quad y = (z > 0).$$

All random variables were sampled from a standard normal distribution. The vectors $\mathbf{a}$ and $\mathbf{d}$ are 16 dimensional vectors that were sampled once to generate the dataset. In $\mathbf{a}$ only the first 4 values have nonzero values to ensure that there are exactly 4 informative features. The values η, ϵ were sampled independently for each example.

We use a least squares model as this problem can be solved linearly. We compare three rankings: the ground truth importance ranking, random ranking and the inverted ground truth ranking (the worst possible estimate of importance). In the left plot of Fig. 2 we can observe that without re-training the worst case estimator is shown to degrade performance relatively quickly. In contrast, ROAR shows no degradation until informative features begin to be removed at 75%. This correctly shows that this estimator has ranked feature importance poorly (ranked uninformative features as most important).

Finally, we consider ROAR performance given a set of variables that are completely redundant. We note that ROAR might not decrease until all of them are removed. To account for this we measure ROAR at different levels of degradation, with the expectation that across this interval we would be able to detect inflection points in performance that would indicate a set of redundant features. If this happens, we believe that it could be detected easily by the sharp decrease as shown in Fig. 2. Now that we have validated ROAR in a controlled setup, we can move on to our large scale experiments.

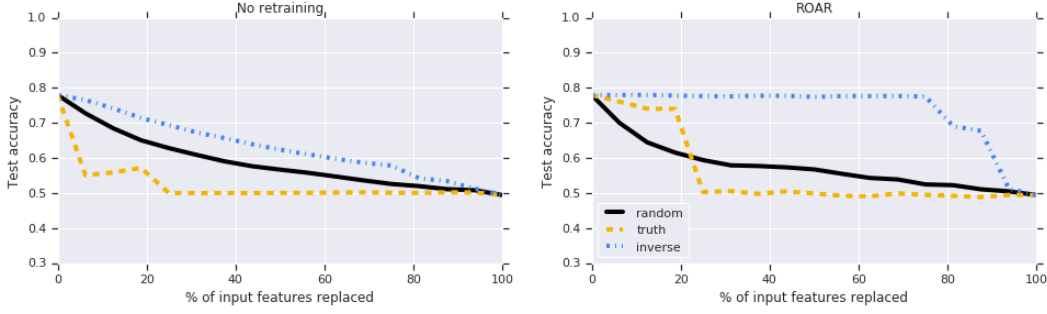


Figure 2: A comparison between not retraining and ROAR on artificial data. In the case where the model is not retrained, test-set accuracy quickly erodes despite the worst case ranking of redundant features as most important. This incorrectly evaluates a completely incorrect feature ranking as being informative. ROAR is far better at identifying this worst case estimator, showing no degradation until the features which are informative are removed at 75%. This plot also shows the limitation of ROAR, an accuracy decrease might not happen until a complete set of fully redundant features is removed. To account for this we measure ROAR at different levels of degradation, with the expectation that across this interval we would be able to control for performance given a set of redundant features.

4 Large scale experiments

4.1 Estimators under consideration

Our evaluation is constrained to a subset of estimators of feature importance. We selected these based on the availability of open source code, consistent guidelines on how to apply them and the ease of implementation given a ResNet-50 architecture [17]. Due to the breadth of the experimental setup it was not possible to include additional methods. However, we welcome the opportunity to consider additional estimators in the future, and in order to make it easy to apply ROAR to additional estimators we have open sourced our code <https://bit.ly/2ttLLZB>. Below, we briefly introduce each of the methods we evaluate.

Base estimators are estimators that compute a single estimate of importance (as opposed to ensemble methods). While we note that guided backprop and integrated gradients are examples of signal and attribution methods respectively, the performance of these estimators should not be considered representative of other methods, which should be evaluated separately.

- **Gradients or Sensitivity heatmaps [34, 7] (GRAD)** are the gradient of the output activation of interest A_n^l with respect to x_i :

$$\mathbf{e} = \frac{\partial A_n^l}{\partial x_i}$$

- **Guided Backprop [37] (GB)** is an example of a signal method that aim to visualize the input patterns that cause the neuron activation A_n^l in higher layers [37, 40, 20]. GB computes this by using a modified backpropagation step that stops the flow of gradients when less than zero at a ReLU gate.
- **Integrated Gradients [38] (IG)** is an example of an attribution method which assign importance to input features by decomposing the output activation A_n^l into contributions from the individual input features [6, 38, 23, 33, 20]. Integrated gradients interpolate a set of estimates for values between a non-informative reference point $\mathbf{x}^0$ to the actual input $\mathbf{x}$. This integral can be approximated by summing a set of k points at small intervals between $\mathbf{x}^0$ and $\mathbf{x}$:

$$\mathbf{e} = (\mathbf{x}_i - \mathbf{x}_i^0) \times \sum_{i=1}^k \frac{\partial f_w(\mathbf{x}^0 + \frac{i}{k}(\mathbf{x} - \mathbf{x}^0))}{\partial \mathbf{x}_i} \times \frac{1}{k}$$

The final estimate $\mathbf{e}$ will depend upon both the choice of k and the reference point $\mathbf{x}^0$. As suggested by [38], we use a black image as the reference point and set k to be 25.

Ensembling methods In addition to the base approaches we also evaluate three ensembling methods for feature importance. For all the ensemble approaches that we describe below (SG, SG-SQ, Var), we average over a set of 15 estimates as suggested by in the original SmoothGrad publication [35].

- **Classic SmoothGrad (SG)** [35] SG averages a set J noisy estimates of feature importance (constructed by injecting a single input with Gaussian noise η independently J times):

$$\mathbf{e} = \sum_{i=1}^J (g_i(\mathbf{x} + \eta, A_n^l))$$

- **SmoothGrad²(SG-SQ)** is an unpublished variant of classic SmoothGrad SG which squares each estimate $\mathbf{e}$ before averaging the estimates:

$$\mathbf{e} = \sum_{i=1}^J (g_i(\mathbf{x} + \eta, A_n^l)^2)$$

Although SG-SQ is not described in the original publication, it is the default open-source implementation of the open source code for SG: <https://bit.ly/2Hpx5ob>.

- **VarGrad (Var)** [2] employs the same methodology as classic SmoothGrad (SG) to construct a set of J noisy estimates. However, VarGrad aggregates the estimates by computing the variance of the noisy set rather than the mean.

$$\mathbf{e} = \text{Var}(g_i(\mathbf{x} + \eta, A_n^l))$$

Control Variants As a control, we compare each estimator to two rankings (a random assignment of importance and a sobel edge filter) that do not depend at all on the model parameters. These controls represent a lower bound in performance that we would expect all interpretability methods to outperform.

- **Random** A random estimator g^R assigns a random binary importance probability $\mathbf{e} \mapsto 0, 1$. This amounts to a binary vector $\mathbf{e} \sim \text{Bernoulli}(1 - t)$ where $(1 - t)$ is the probability of $e_i = 1$. The formulation of g^R does not depend on either the model parameters or the input image (beyond the number of pixels in the image).
- **Sobel Edge Filter** convolves a hard-coded, separable, integer filter over an image to produce a mask of derivatives that emphasizes the edges in an image. A sobel mask treated as a ranking $\mathbf{e}$ will assign a high score to areas of the image with a high gradient (likely edges).

4.2 Experimental setup

We use a ResNet-50 model for both generating the feature importance estimates and subsequent training on the modified inputs. ResNet-50 was chosen because of the public code implementations (in both PyTorch [15] and Tensorflow [1]) and because it can be trained to give near to state of art performance in a reasonable amount of time [14].

For all train and validation images in the dataset we first apply test time pre-processing as used by Goyal et al. [14]. We evaluate ROAR on three open source image datasets: ImageNet, Birdsnap and Food 101. For each dataset and estimator, we generate new train and test sets that each correspond to a different fraction of feature modification $t = [0, 10, 30, 50, 70, 90]$. We evaluate 18 estimators in total (this includes the base estimators, a set of ensemble approaches wrapped around each base and finally a set of squared estimates). In total, we generate 540 large-scale modified image datasets in order to consider all experiment variants (180 new test/train for each original dataset).

We independently train 5 ResNet-50 models from random initialization on each of these modified dataset and report test accuracy as the average of these 5 runs. In the base implementation, the ResNet-50 trained on an unmodified ImageNet dataset achieves a mean accuracy of 76.68%. This is comparable to the performance reported by [14]. On Birdsnap and Food 101, our unmodified datasets achieve 66.65% and 84.54% respectively (average of 10 independent runs). This baseline performance is comparable to that reported by Kornblith et al. [21].

4.3 Experimental results

4.3.1 Evaluating the random ranking

Comparing estimators to the random ranking allows us to answer the question: *is the estimate of importance more accurate than a random guess?* It is firstly worthwhile noting that model performance is remarkably robust to random modification. After replacing a large portion of all inputs with a constant value, the model not only trains but still retains most of the original predictive power. For example, on ImageNet, when only 10% of all features are retained, the trained model still attains 63.53% accuracy (relative to unmodified baseline of 76.68%). The ability of the model to extract a meaningful representation from a small random fraction of inputs suggests a case where many inputs are likely redundant. The nature of our input—an image where correlations between pixels are expected – provides one possible readons for redundancy.

The results for our random baseline provides additional support for the need to re-train. We can compare random ranking on ROAR vs. a traditional deletion metric [30], i.e. the setting where we do not retain. These results are given in Fig. 3. Without retraining, a random modification of 90% degrades accuracy to 0.5% for the model that was not retrained. Keep in mind that on clean data we achieve 76.68% accuracy. This large discrepancy illustrates that without retraining the model, it is not possible to decouple the performance of the ranking from the degradation caused by the modification itself.

4.3.2 Evaluating Base Estimators

Now that we have established the baselines, we can start evaluating the base estimators: GB, IG, GRAD. Surprisingly, the left inset of Fig. 4 shows that these estimators consistently perform worse than the random assignment of feature importance across all datasets and for all thresholds $t = [0.1, 0.3, 0.5, 0.7, 0.9]$. Furthermore, our estimators fall further behind the accuracy of random guess as a larger fraction t of inputs is modified. The gap is widest when $t = 0.9$.

Our base estimators also do not compare favorably to the performance of a sobel edge filter SOBEL. Both the sobel filter and the random ranking have formulations that are entirely independent of the model parameters. All the base estimators that we consider have formulations that depend upon the trained model weights, and thus we would expect them to have a clear advantage in outperforming the control variants. However, across all datasets and thresholds t , the base estimators GB, IG, GRAD perform on par or worse than SOBEL.

Base estimators perform within a very narrow range. Despite the very different formulations of base estimators that we consider, the difference between the performance of the base estimators is in a strikingly narrow range. For example, as can be seen in the left column of Fig. 4, for Birdsnap, the difference in accuracy between the best and worst base estimator at $t = 90\%$ is only 4.22%. This range remains narrow for both Food101 and ImageNet, with a gap of 5.17% and 3.62 respectively. Our base estimator results are remarkably consistent results across datasets, methods and for all fractions of t considered. The variance is very low across independent runs for all datasets and estimators. The maximum variance observed for ImageNet was a variance of 1.32% using SG-SQ-GRAD at 70% of inputs removed. On Birdsnap the highest variance was 0.12% using VAR-GRAD at 90% removed. For food101 it was 1.52% using SG-SQ-GRAD at 70% removed.

Finally, we compare performance of the base estimators using ROAR re-training vs. a traditional deletion metric [30], again the setting where we do not retain. In Fig. 3 we see a behavior for the base estimators on all datasets that is similar to the behavior of the inverse (worst possible) ranking on the toy data in Fig. 2. The base estimators appear to be working when we do not retrain, but they are clearly not better than the random baseline when evaluated using ROAR. This provides additional support for the need to re-train.

4.3.3 Evaluating Ensemble Approaches

Since the base estimators do not appear to perform well, we move on to ensemble estimators. Ensemble approaches inevitably carry a higher computational approach, as the methodology requires the aggregation of a set of individual estimates. However, these methods are often preferred by humans because they appear to produce “less noisy” explanations. However, there is limited theoretical understanding of what these methods are actually doing or how this is related to the accuracy of the

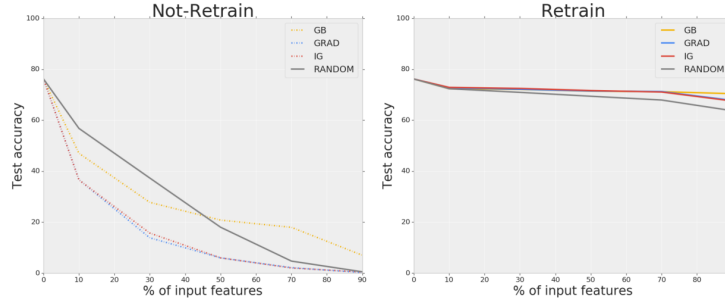


Figure 3: On the left we evaluate three base estimators and the random baseline without retraining. All of the methods appear to reduce accuracy at quite a high rate. On the right, we see, using ROAR, that after re-training most of the information is actually still present. It is also striking that in this case the base estimators perform worse than the random baseline.

Only Certain Ensemble Approaches Benefit Performance

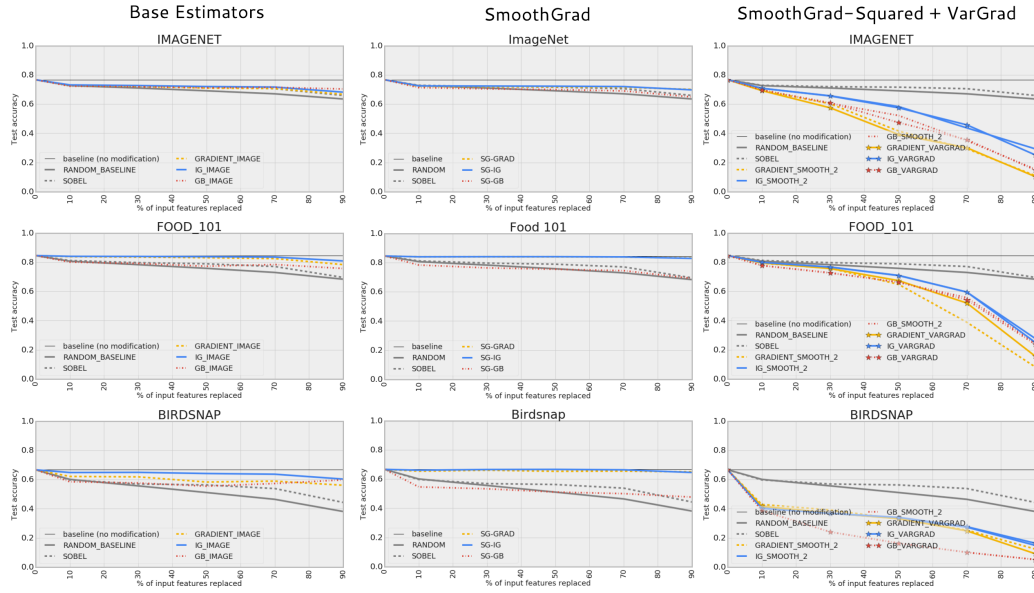


Figure 4: **Left:** Grad (GRAD), Integrated Gradients (IG) and Guided Backprop (GB) perform worse than a random assignment of feature importance. **Middle:** SmoothGrad (SG) is less accurate than a random assignment of importance and often worse than a single estimate (in the case of raw gradients SG-Grad and Integrated Gradients SG-IG). **Right:** SmoothGrad Squared (SG-SQ) and VarGrad (VAR) produce a dramatic improvement in approximate accuracy and far outperform the other methods in all datasets considered, regardless of the underlying estimator.

explanation. We evaluate ensemble estimators and produce results that are remarkably consistent results across datasets, methods and for all fractions of t considered.

Classic Smoothgrad is less accurate or on par with a single estimate. In the middle column in Fig. 4 we evaluate Classic SmoothGrad (SG). It average 15 estimates computed according to an underlying base method (GRAD, IG or GB). However, despite averaging SG degrades test-set accuracy still less than a random guess. In addition, for GRAD and IG SmoothGrad performs worse than a single estimate.

SmoothGrad-Squared and VarGrad produce large gains in accuracy. In the right inset of Fig. 4, we show that both VarGrad (VAR) and SmoothGrad-Squared (SG-SQ) far outperform the two control variants. In addition, for all the interpretability methods we consider, VAR or SG-SQ far outperform the approximate accuracy of a single estimate. However, while VAR and SG-SQ benefit the accuracy of all base estimators, the overall ranking of estimator performance differs by dataset. For ImageNet and Food101, the best performing estimators are VAR or SG-SQ when wrapped around GRAD. However, for the Birdsnap dataset, the most approximately accurate estimates are these ensemble approaches wrapped around GB. This suggests that while the VAR and SG-SQ consistently improve performance, the choice of the best underlying estimator may vary by task.

Now, why do both of these methods work so well? First, these methods are highly similar. If the average (squared) gradient over the noisy samples is zero then VAR and SG-SQ reduce to the same method. For many images it appears that the mean gradient is much smaller than the mean squared gradient. This implies that the final output should be similar. Qualitatively this seems to be the case as well. In Fig. 1 we observe that both methods appear to remove whole objects. The other methods removed inputs that are less concentrated but spread more widely over the image. It is important to note that these methods were not forced to behave as such. It is emergent behavior. Understanding why this happens and why this is beneficial should be the focus of future work.

Squaring estimates The final question we consider is why SmoothGrad-Squared SG-SQ dramatically improves upon the performance of SmoothGrad SG despite little difference in formulation. The only difference between the two estimates is that SG-SQ squares the estimates before averaging. We consider the effect of only squaring estimates (no ensembling). We find that while squaring improves the accuracy of all estimators, the transformation does not adequately explain the large gains that we observe when applying VAR or SG-SQ. When base estimators are squared they slightly outperform the random baseline (all results included in the supplementary materials).

5 Conclusion and Future Work

In this work, we propose ROAR to evaluate the quality of input feature importance estimators. Surprisingly, we find that the commonly used base estimators, Gradients, Integrated Gradients and Guided BackProp are worse or on par with a random assignment of importance. Furthermore, certain ensemble approaches such as SmoothGrad are far more computationally intensive but do not improve upon a single estimate (and in some cases are worse). However, we do find that VarGrad and SmoothGrad-Squared strongly improve the quality of these methods and far outperform a random guess. While the low effectiveness of many methods could be seen as a negative result, we view the remarkable effectiveness of SmoothGrad-Squared and VarGrad as important progress within the community. Our findings are particularly pertinent for sensitive domains where the accuracy of a explanation of model behavior is paramount. While we venture some initial consideration of why certain ensemble methods far outperform other estimator, the divergence in performance between the ensemble estimators is an important direction of future research.

Acknowledgments

We thank Gabriel Bender, Kevin Swersky, Andrew Ross, Douglas Eck, Jonas Kemp, Melissa Fabros, Julius Abedayo, Simon Kornblith, Prajit Ramachandran, Niru Maheswaranathan, Gamaleldin Elsayed, Hugo Larochelle, Varun Vasudevan for their thoughtful feedback on earlier iterations of this work. In addition, thanks to Sally Jesmonth, Dan Nanas and Alexander Popper for institutional support and encouragement.

References

- [1] Martín Abadi, Ashish Agarwal, Paul Barham, Eugene Brevdo, Zhifeng Chen, Craig Citro, Greg S. Corrado, Andy Davis, Jeffrey Dean, Matthieu Devin, Sanjay Ghemawat, Ian Goodfellow, Andrew Harp, Geoffrey Irving, Michael Isard, Yangqing Jia, Rafal Jozefowicz, Lukasz Kaiser, Manjunath Kudlur, Josh Levenberg, Dandelion Mané, Rajat Monga, Sherry Moore, Derek Murray, Chris Olah, Mike Schuster, Jonathon Shlens, Benoit Steiner, Ilya Sutskever, Kunal Talwar, Paul Tucker, Vincent Vanhoucke, Vijay Vasudevan, Fernanda Viégas, Oriol Vinyals, Pete Warden, Martin Wattenberg, Martin Wicke, Yuan Yu, and Xiaoqiang Zheng. TensorFlow:

- Large-scale machine learning on heterogeneous systems, January 2015. Software available from tensorflow.org.
- [2] Julius Adebayo, Justin Gilmer, Ian Goodfellow, and Been Kim. Local explanation methods for deep neural networks lack sensitivity to parameter values. *ICLR Workshop*, 2018.
 - [3] Julius Adebayo, Justin Gilmer, Michael Muelly, Ian J. Goodfellow, Moritz Hardt, and Been Kim. Sanity checks for saliency maps. In *NeurIPS*, 2019.
 - [4] M. Ancona, E. Ceolini, C. Öztireli, and M. Gross. Towards better understanding of gradient-based attribution methods for Deep Neural Networks. *ArXiv e-prints*, November 2017.
 - [5] Lei Jimmy Ba and Rich Caruana. Do deep nets really need to be deep? In *Proceedings of the 27th International Conference on Neural Information Processing Systems - Volume 2, NIPS’14*, pages 2654–2662, Cambridge, MA, USA, 2014. MIT Press.
 - [6] Sebastian Bach, Alexander Binder, Grégoire Montavon, Frederick Klauschen, Klaus-Robert Müller, and Wojciech Samek. On pixel-wise explanations for non-linear classifier decisions by layer-wise relevance propagation. *PloS one*, 10(7):e0130140, 2015.
 - [7] David Baehrens, Timon Schroeter, Stefan Harmeling, Motoaki Kawanabe, Katja Hansen, and Klaus-Robert Müller. How to explain individual classification decisions. *Journal of Machine Learning Research*, 11(Jun):1803–1831, 2010.
 - [8] Thomas Berg, Jiongxin Liu, Seung Woo Lee, Michelle L. Alexander, David W. Jacobs, and Peter N. Belhumeur. Birdsnap: Large-scale fine-grained visual categorization of birds. *2014 IEEE Conference on Computer Vision and Pattern Recognition*, pages 2019–2026, 2014.
 - [9] Lukas Bossard, Matthieu Guillaumin, and Luc Van Gool. Food-101 – mining discriminative components with random forests. In *European Conference on Computer Vision*, 2014.
 - [10] P. Dabkowski and Y. Gal. Real Time Image Saliency for Black Box Classifiers. *ArXiv e-prints*, May 2017.
 - [11] J. Deng, W. Dong, R. Socher, L.-J. Li, K. Li, and L. Fei-Fei. ImageNet: A Large-Scale Hierarchical Image Database. In *CVPR09*, 2009.
 - [12] Ruth C. Fong and Andrea Vedaldi. Interpretable explanations of black boxes by meaningful perturbation. In *ICCV*, pages 3449–3457. IEEE Computer Society, 2017.
 - [13] N. Frosst and G. Hinton. Distilling a Neural Network Into a Soft Decision Tree. *ArXiv e-prints*, November 2017.
 - [14] P. Goyal, P. Dollár, R. Girshick, P. Noordhuis, L. Wesolowski, A. Kyrola, A. Tulloch, Y. Jia, and K. He. Accurate, Large Minibatch SGD: Training ImageNet in 1 Hour. *ArXiv e-prints*, June 2017.
 - [15] S. Gross and M. Wilber. Training and investigating Residual Nets. <https://github.com/facebook/fb.resnet.torch>, January 2017.
 - [16] Isabelle Guyon, Jason Weston, Stephen Barnhill, and Vladimir Vapnik. Gene selection for cancer classification using support vector machines. *Machine Learning*, 46(1):389–422, Jan 2002.
 - [17] K. He, X. Zhang, S. Ren, and J. Sun. Deep Residual Learning for Image Recognition. *ArXiv e-prints*, December 2015.
 - [18] Been Kim, Martin Wattenberg, Justin Gilmer, Carrie Cai, James Wexler, Fernanda Viegas, and Rory sayres. Interpretability beyond feature attribution: Quantitative testing with concept activation vectors (TCAV). In Jennifer Dy and Andreas Krause, editors, *Proceedings of the 35th International Conference on Machine Learning*, volume 80 of *Proceedings of Machine Learning Research*, pages 2668–2677, Stockholmsmässan, Stockholm Sweden, 10–15 Jul 2018. PMLR.

- [19] P.-J. Kindermans, S. Hooker, J. Adebayo, M. Alber, K. T. Schütt, S. Dähne, D. Erhan, and B. Kim. The (Un)reliability of saliency methods. *ArXiv e-prints*, November 2017.
- [20] Pieter-Jan Kindermans, Kristof T. Schütt, Maximilian Alber, Klaus-Robert Müller, Dumitru Erhan, Been Kim, and Sven Dähne. Learning how to explain neural networks: Patternnet and patternattribution. In *International Conference on Learning Representations*, 2018.
- [21] Simon Kornblith, Jonathon Shlens, and Quoc V. Le. Do Better ImageNet Models Transfer Better? *arXiv e-prints*, page arXiv:1805.08974, May 2018.
- [22] Isaac Lage, Andrew Slavin Ross, Been Kim, Samuel J. Gershman, and Finale Doshi-Velez. Human-in-the-Loop Interpretability Prior. *arXiv e-prints*, page arXiv:1805.11571, May 2018.
- [23] Grégoire Montavon, Sebastian Lapuschkin, Alexander Binder, Wojciech Samek, and Klaus-Robert Müller. Explaining nonlinear classification decisions with deep taylor decomposition. *Pattern Recognition*, 65:211–222, 2017.
- [24] A. S. Morcos, D. G. T. Barrett, N. C. Rabinowitz, and M. Botvinick. On the importance of single directions for generalization. *ArXiv e-prints*, March 2018.
- [25] Chris Olah, Alexander Mordvintsev, and Ludwig Schubert. Feature visualization. *Distill*, 2017. <https://distill.pub/2017/feature-visualization>.
- [26] Vitali Petsiuk, Abir Das, and Kate Saenko. RISE: randomized input sampling for explanation of black-box models. *CoRR*, abs/1806.07421, 2018.
- [27] Maithra Raghu, Justin Gilmer, Jason Yosinski, and Jascha Sohl-Dickstein. SVCCA: singular vector canonical correlation analysis for deep learning dynamics and interpretability. In *NIPS*, pages 6078–6087, 2017.
- [28] Andrew Slavin Ross and Finale Doshi-Velez. Improving the adversarial robustness and interpretability of deep neural networks by regularizing their input gradients. *CoRR*, abs/1711.09404, 2017.
- [29] Andrew Slavin Ross, Michael C. Hughes, and Finale Doshi-Velez. Right for the right reasons: Training differentiable models by constraining their explanations. In *Proceedings of the Twenty-Sixth International Joint Conference on Artificial Intelligence, IJCAI 2017, Melbourne, Australia, August 19-25, 2017*, pages 2662–2670, 2017.
- [30] W. Samek, A. Binder, G. Montavon, S. Lapuschkin, and K. R. Müller. Evaluating the Visualization of What a Deep Neural Network Has Learned. *IEEE Transactions on Neural Networks and Learning Systems*, 28(11):2660–2673, Nov 2017.
- [31] Ramprasaath R. Selvaraju, Michael Cogswell, Abhishek Das, Ramakrishna Vedantam, Devi Parikh, and Dhruv Batra. Grad-cam: Visual explanations from deep networks via gradient-based localization. In *The IEEE International Conference on Computer Vision (ICCV)*, Oct 2017.
- [32] A. Shrikumar, P. Greenside, and A. Kundaje. Learning Important Features Through Propagating Activation Differences. *ArXiv e-prints*, April 2017.
- [33] A. Shrikumar, P. Greenside, A. Shcherbina, and A. Kundaje. Not Just a Black Box: Learning Important Features Through Propagating Activation Differences. *ArXiv e-prints*, May 2016.
- [34] Karen Simonyan and Andrew Zisserman. Very deep convolutional networks for large-scale image recognition. In *ICLR*, 2015.
- [35] Daniel Smilkov, Nikhil Thorat, Been Kim, Fernanda Viégas, and Martin Wattenberg. Smoothgrad: removing noise by adding noise. *arXiv preprint arXiv:1706.03825*, 2017.
- [36] Irwin Sobel. An isotropic 3x3 image gradient operator. *Presentation at Stanford A.I. Project 1968*, 02 2014.
- [37] Jost Tobias Springenberg, Alexey Dosovitskiy, Thomas Brox, and Martin Riedmiller. Striving for simplicity: The all convolutional net. In *ICLR*, 2015.

- [38] Mukund Sundararajan, Ankur Taly, and Qiqi Yan. Axiomatic attribution for deep networks. *arXiv preprint arXiv:1703.01365*, 2017.
- [39] M. Wu, M. C. Hughes, S. Parbhoo, M. Zazzi, V. Roth, and F. Doshi-Velez. Beyond Sparsity: Tree Regularization of Deep Models for Interpretability. *ArXiv e-prints*, November 2017.
- [40] Matthew D Zeiler and Rob Fergus. Visualizing and understanding convolutional networks. In *European Conference on Computer Vision*, pages 818–833. Springer, 2014.
- [41] C. Zhang, S. Bengio, M. Hardt, B. Recht, and O. Vinyals. Understanding deep learning requires rethinking generalization. *ArXiv e-prints*, November 2016.
- [42] B. Zhou, A. Khosla, A. Lapedriza, A. Oliva, and A. Torralba. Object Detectors Emerge in Deep Scene CNNs. *ArXiv e-prints*, 2014.
- [43] Bolei Zhou, Yiyu Sun, David Bau, and Antonio Torralba. Revisiting the importance of individual units in cnns via ablation. *CoRR*, abs/1806.02891, 2018.
- [44] Luisa M Zintgraf, Taco S Cohen, Tameem Adel, and Max Welling. Visualizing deep neural network decisions: Prediction difference analysis. In *ICLR*, 2017.

6 Supplementary Charts and Experiments

We include supplementary experiments and additional details about our training procedure, the estimators we evaluate, the image modification process and test-set accuracy below. In addition, as can be seen in Fig. 5, we also consider the scenario where pixels are kept according to importance rather than removed.

6.1 Training Procedure

We carefully tuned the hyperparameters of each dataset ImageNet, Birdsnap and Food 101 separately. We find that the Birdsnap and Food 101 converge within the same amount of training steps and a larger learning rate than ImageNet. These are detailed in Table. 6.2. These hyper parameters, along with the mean accuracy reported on the unmodified dataset, are used consistently across all estimators. ImageNet dataset achieves a mean accuracy of 76.68%. This is comparable to the performance reported by [14]. On Birdsnap and Food 101, our unmodified datasets achieve 66.65% and 84.54% respectively. The baseline test-set accuracy for Food101 or Birdsnap is comparable to that reported by [21]. In Table. 2, we include the test-set performance for each experiment variant that we consider. The test-set accuracy reported is the average of 5 independent runs.

6.2 Generation of New Dataset

Dataset	Top1Accuracy	Train Size	Test Size	Learning Rate	Training Steps
Birdsnap	66.65	47,386	2,443	1.0	20,000
Food_101	84.54	75,750	25,250	0.7	20,000
ImageNet	76.68	1,281,167	50,000	0.1	32,000

Table 1: The training procedure was carefully finetuned for each dataset. These hyperparameters are consistently used across all experiment variants. The baseline accuracy of each unmodified data set is reported as the average of 10 independent runs.

We evaluate ROAR on three open source image datasets: ImageNet, Birdsnap and Food 101. For each dataset and estimator, we generate 10 new train and test sets that each correspond to a different fraction of feature modification $t = [0.1, 0.3, 0.5, 0.7, 0.9]$ and whether the most important pixels are removed or kept. This requires first generating a ranking of input importance for each input image according to each estimator. All of the estimators that we consider evaluate feature importance post-training. Thus, we generate the rankings according to each interpretability method using a stored checkpoint for each dataset.

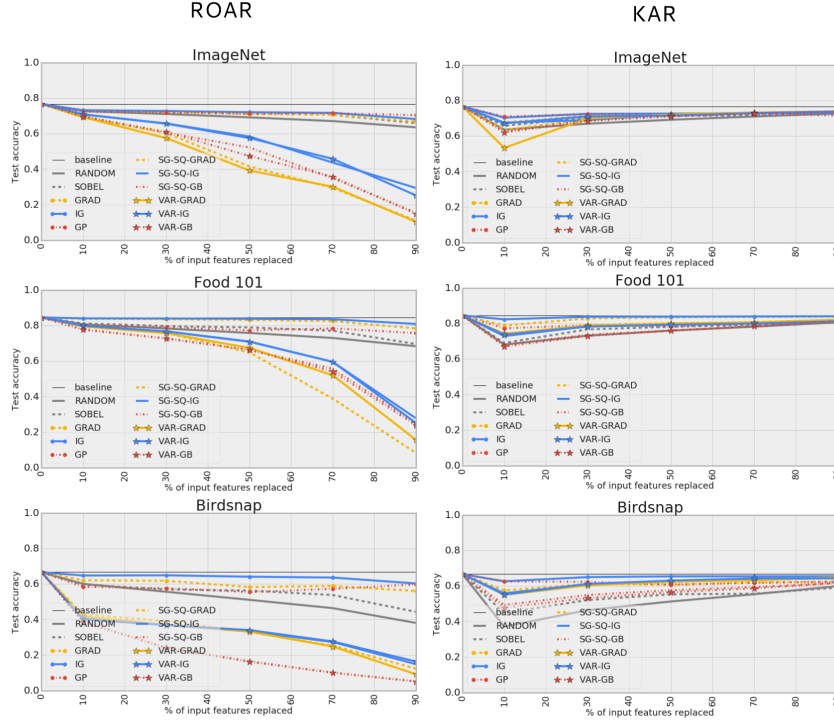


Figure 5: Evaluation of all estimators according to Keep and Retrain KAR vs. ROAR. **Left inset:** For KAR, **Keep And Retrain**, we keep a fraction of features estimated to be most important and replace the remaining features with a constant mean value. The most accurate estimator is the one that preserves model performance the most for a given fraction of inputs removed (the highest test-set accuracy). **Right inset:** For **ROAR, Remove And Retrain** we remove features by replacing a fraction of the inputs estimated to be most important according to each estimator with a constant mean value. The most accurate estimator is the one that degrades model performance the most for a given fraction of inputs removed. Inputs modified according to KAR result in a very narrow range of model accuracy. ROAR is a more discriminative benchmark, which suggests that retaining performance when the most important pixels are removed (rather than retained) is a harder task.

We use the ranking produced by the interpretability method to modify each image in the dataset (both train and test). We rank each estimate, e into an ordered set $\{e_i^o\}_{i=1}^N$. For the top t fraction of this ordered set, we replace the corresponding pixels in the raw image with a per channel mean. Fig. 7 and Fig. 8 show an example of the type of modification applied to each image in the dataset for Birdsnap and Food 101 respectively. In the paper itself, we show an example of a single image from each ImageNet modification.

We evaluate 18 estimators in total (this includes the base estimators, a set of ensemble approaches wrapped around each base and finally a set of squared estimates). In total, we generate 540 large-scale modified image datasets in order to consider all experiment variants (180 new test/train for each original dataset).

6.3 Evaluating Keeping Rather Than Removing Information

In addition to ROAR, as can be seen in Fig. 5, we evaluate the opposite approach of **KAR, Keep And Retrain**. While ROAR removes features by replacing a fraction of inputs estimated to be most important, KAR preserves the inputs considered to be most important. Since we keep the important information rather than remove it, *minimizing* degradation to test-set accuracy is desirable.

In the right inset chart of Fig. 5 we plot KAR on the same curve as ROAR to enable a more intuitive comparison between the benchmarks. The comparison suggests that KAR appears to be

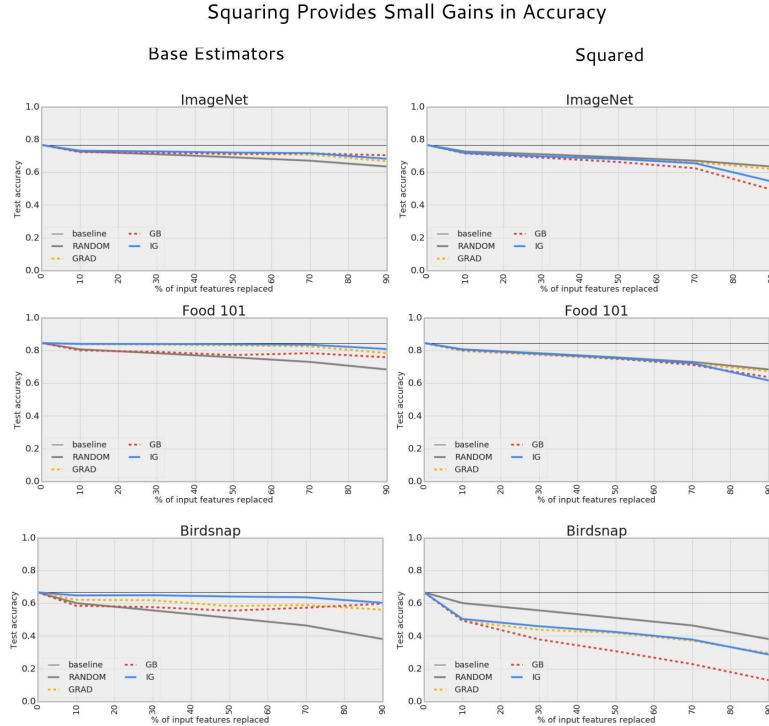


Figure 6: Certain transformations of the estimate can substantially improve accuracy of all estimators. Squaring alone provides small gains to the accuracy of all estimators, and is slightly better than a random guess. **Left inset:** The three base estimators that we consider (Gradients (GRAD), Integrated Gradients (IG) and Guided Backprop (GB)) perform worse than a random assignment of feature importance. At all fractions considered, a random assignment of importance degrades performance more than removing the pixels estimated to be most important by base methods. **Right inset:** Average test-set accuracy across 5 independent iterations for estimates that are squared before ranking and subsequent removal. When squared, base estimators perform slightly better than a random guess. However, this does not compare to the gains in accuracy of averaging a set of noisy estimates that are squared (SmoothGrad-Squared)

a poor discriminator between estimators. The x-axis indicates the fraction of features that are preserved/removed for KAR/ROAR respectively.

We find that KAR is a far weaker discriminator of performance; all base estimators and the ensemble variants perform in a similar range to each other. These findings suggest that the task of identifying features to preserve is an easier benchmark to fulfill than accurately identifying a fraction of input that will cause the maximum damage to the model performance.

6.4 Squaring Alone Slightly Improves the Performance of All Base Variants

The surprising performance of SmoothGrad-Squared (SG-SQ) deserves further investigation; why is averaging a set of squared noisy estimates so effective at improving the accuracy of the ranking? To disentangle whether both squaring and then averaging are required, we explore whether we achieve similar performance gains by **only** squaring the estimate.

Squaring of a single estimate, with no ensembling, benefits the accuracy of all estimators that we considered. In the right inset chart of Fig. 6, we can see that squared estimates perform better than the raw estimate. When squared, an estimate gains slightly more accuracy than a random ranking of input features. In particular, squaring benefits GB; at $t = .9$ performance of SQ-GB relative to GB improves by $8.43\% \pm 0.97$.



Figure 7: A single example from each dataset generated from modifying Food 101 according to both ROAR and KAR. We show the modification for base estimators (Integrated Gradients (IG), Guided Backprop (GB), Gradient Heatmap (GRAD) and derivative ensemble approaches - SmoothGrad, (SG-GRAD, SG-IG, SG-GB), SmoothGrad-Squared (SG-SQ-GRAD, SG-SQ-IG, SG-SQ-GB) and VarGrad (VAR-GRAD, VAR-IG, VAR-GB). In addition, we consider two control variants a random baseline, a sobel edge filter.

Squaring is an equivalent transformation to taking the absolute value of the estimate before ranking all inputs. After squaring, negative estimates become positive, and the ranking then only depends upon the magnitude and not the direction of the estimate. The benefits gained by squaring furthers our understanding of how the direction of GB, IG and GRAD values should be treated. For all these estimators, estimates are very much a reflection of the weights of the network. The magnitude may be far more telling of feature importance than direction; a negative signal may be just as important as positive contributions towards a model’s prediction. While squaring improves the accuracy of all estimators, the transformation does not explain the large gains in accuracy that we observe when we average a set of noisy squared estimates.

6.5 Limitations on the use of ROAR

In this work we propose ROAR as a method for estimating feature importance in deep neural networks. However, we do note that ROAR is not suitable for certain algorithms such as decision stump $Y=(A \text{ or } D)$ where there is also feature redundancy. For these algorithms, in order to use ROAR correctly feature importance must be recomputed after each re-training step. This is because a decision stump

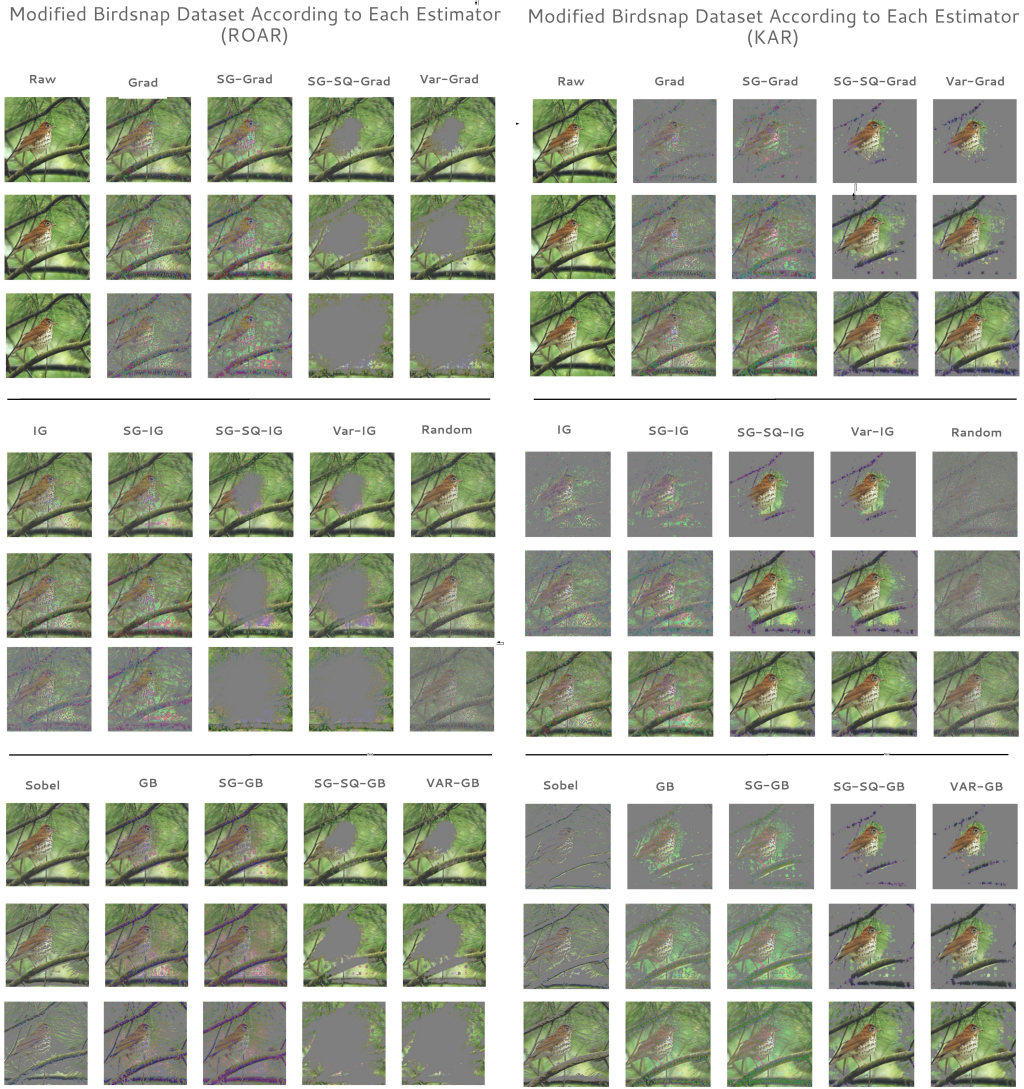


Figure 8: A single example from each dataset generated from modifying Imagenet according to the ROAR and KAR. We show the modification for base estimators (Integrated Gradients (IG), Guided Backprop (GB), Gradient Heatmap (GRAD) and derivative ensemble approaches - SmoothGrad, (SG-GRAD, SG-IG, SG-GB), SmoothGrad-Squared (SG-SQ-GRAD, SG-SQ-IG, SG-SQ-GB) and VarGrad (VAR-GRAD, VAR-IG, VAR-GB). In addition, we consider two control variants a random baseline, a sobel edge filter.

ignores a subset of input features at inference time which means it is possible for a random estimator to appear to perform better than the best possible estimator. For the class of models evaluated in the paper (linear models, multi-layer perceptrons and deep neural networks) as well as any model that allows all features to contribute to a prediction at test time ROAR remains valid. To make ROAR valid for decisions stumps, one can re-compute the feature importance after each re-training step. The scale of our experiments preclude this, and our experiments show that it is not necessary for deep neural networks (DNNs).

	Threshold	Keep					Remove				
		10.0	30.0	50.0	70.0	90.0	10.0	30.0	50.0	70.0	90.0
Birdsnap	Random	37.24	46.41	51.29	55.38	59.92	60.11	55.65	51.10	46.45	38.12
	Sobel	44.81	52.11	55.36	55.69	59.08	59.73	56.94	56.30	53.82	44.33
	GRAD	57.51	61.10	60.79	61.96	62.49	62.12	61.82	58.29	58.91	56.08
	IG	62.64	65.02	65.42	65.46	65.50	64.79	64.91	64.12	63.64	60.30
	GP	62.59	62.35	60.76	61.78	62.44	58.47	57.64	55.47	57.28	59.76
	SG-GRAD	64.64	65.87	65.32	65.49	65.78	65.44	66.08	65.33	65.44	65.02
	SG-IG	65.36	66.45	66.38	66.37	66.35	66.11	66.56	66.65	66.37	64.54
	SG-GB	52.86	56.44	58.32	59.20	60.35	54.67	53.37	51.13	50.07	47.71
	SG-SQ-GRAD	55.32	60.79	62.13	63.63	64.99	42.88	39.14	32.98	25.34	12.40
	SG-SQ-IG	55.89	61.02	62.68	63.63	64.43	40.85	36.94	33.37	27.38	14.93
	SG-SQ-GB	49.32	54.94	57.62	59.41	61.66	38.80	24.09	16.54	10.11	5.21
	VAR-GRAD	55.03	60.36	62.59	63.16	64.85	41.71	37.04	33.24	24.84	9.23
	VAR-IG	55.21	61.22	63.04	64.29	64.31	40.21	36.85	34.09	27.71	16.43
	VAR-GB	47.76	53.27	56.53	58.68	61.69	38.63	24.12	16.29	10.16	5.20
Food_101	Random	68.13	73.15	76.00	78.21	80.61	80.66	78.30	75.80	72.98	68.37
	Sobel	69.08	76.70	78.16	79.30	80.90	81.17	79.69	78.91	77.06	69.58
	GRAD	78.82	82.89	83.43	83.68	83.88	83.79	83.50	83.09	82.48	78.36
	IG	82.35	83.80	83.90	83.99	84.07	84.01	83.95	83.78	83.52	80.87
	GP	77.31	79.00	78.33	79.86	81.16	80.06	79.12	77.25	78.43	75.69
	SG-GRAD	83.30	83.87	84.01	84.05	83.96	83.97	84.00	83.97	83.83	83.14
	SG-IG	83.27	83.91	84.06	84.05	83.96	83.98	84.04	84.05	83.90	82.90
	SG-GB	71.44	75.96	77.26	78.65	80.12	78.35	76.39	75.44	74.50	69.19
	SG-SQ-GRAD	73.05	79.20	80.18	80.80	82.13	79.29	75.83	64.83	38.88	8.34
	SG-SQ-IG	72.93	78.36	79.33	80.02	81.30	79.73	76.73	70.98	59.55	27.81
	SG-SQ-GB	68.10	73.69	76.02	78.51	81.22	77.68	72.81	66.24	55.73	24.95
	VAR-GRAD	74.24	78.86	79.97	80.61	82.10	79.55	75.67	67.40	52.05	15.69
	VAR-IG	73.65	78.28	79.31	79.99	81.23	79.87	76.60	70.85	59.57	25.15
	VAR-GB	67.08	73.00	76.01	78.54	81.44	77.76	72.56	66.36	54.18	23.88
ImageNet	Random	63.60	66.98	69.18	71.03	72.69	72.65	71.02	69.13	67.06	63.53
	Sobel	65.79	70.40	71.40	71.60	72.65	72.89	71.94	71.61	70.56	65.94
	GRAD	67.63	71.45	72.02	72.85	73.46	72.94	72.22	70.97	70.72	66.75
	IG	70.38	72.51	72.66	72.88	73.32	73.17	72.72	72.03	71.68	68.20
	GP	71.03	72.45	72.28	72.69	71.56	72.29	71.91	71.18	71.48	70.38
	SG-GRAD	70.47	71.94	72.14	72.35	72.44	72.08	71.94	71.77	71.51	70.10
	SG-IG	70.98	72.30	72.49	72.60	72.67	72.49	72.39	72.26	72.02	69.77
	SG-GB	66.97	70.68	71.52	71.86	72.57	71.28	70.45	69.98	69.02	64.93
	SG-SQ-GRAD	63.25	69.79	72.20	73.18	73.96	69.35	60.28	41.55	29.45	11.09
	SG-SQ-IG	67.55	68.96	72.24	73.09	73.80	70.76	65.71	58.34	43.71	29.41
	SG-SQ-GB	62.42	68.96	71.17	72.72	73.77	69.74	60.56	52.21	34.98	15.53
	VAR-GRAD	53.38	69.86	72.15	73.22	73.92	69.24	57.48	39.23	30.13	10.41
	VAR-IG	67.17	71.07	71.48	72.93	73.87	70.87	65.56	57.49	45.80	25.25
	VAR-GB	62.09	68.51	71.09	72.59	73.85	69.67	60.94	47.39	35.68	14.93

Table 2: Average test-set accuracy across 5 independent runs for all estimators and datasets considered. ROAR requires removing a fraction of pixels estimated to be most important. KAR differs in that the pixels estimated to be most important are kept rather than removed. The fraction removed/kept is indicated by the threshold. The estimators we report results for are base estimators (Integrated Gradients (IG), Guided Backprop (GB), Gradient Heatmap (GRAD) and derivative ensemble approaches - SmoothGrad, (SG-GRAD, SG-IG, SG-GB), SmoothGrad-Squared (SG-SQ-GRAD, SG-SQ-IG, SG-SQ-GB) and VarGrad (VAR-GRAD, VAR-IG, VAR-GB).