
Supplementary material:

Gradient strikes back: How filtering out high frequencies improves explanations

Anonymous Author(s)

Affiliation

Address

email

1 A Proofs

2 In this section, we build on the assumption – that we empirically showed in the paper – that the
 3 gradient is noisy and demonstrate, through a Fourier perspective, that **FORGrad** method effectively
 4 recovers the true gradient. We then propose a convergence bound for SmoothGrad, showing that it
 5 also recovers the true gradient at the cost of multiple samplings, and provide convergence bounds as
 6 well.

As a recall, we still consider a predictor $f(\cdot)$ that maps datum from an input space $\mathcal{X} \subseteq \mathbb{R}^{W \times H}$ (with W, H being positive integers) to an output space $\mathcal{Y} \subseteq \mathbb{R}$. Moreover, $\mathcal{F}$ and $\mathcal{F}^{-1}$ still denote the Discrete Fourier Transform (DFT) on $\mathcal{X}$ and its inverse. We assume that f is L -lipschitz. We recall that a function f is said L -lipschitz $f \in Lip(\mathcal{X})$ if and only if $\forall (x, z) \in \mathcal{X}^2, \|f(x) - f(z)\| \leq L\|x - z\|$, with $L \in \mathbb{R}$. Finally, we define $K^\sigma \in \{0, 1\}^{W \times H}$ as the binary mask parametrized by σ that we used to filter high frequency, where each element $K_{(i,j)}^\sigma$ is determined as follows:

$$K_{(i,j)}^\sigma = \begin{cases} 1, & \text{if } |i - \frac{W}{2}| \leq \sigma \text{ and } |j - \frac{H}{2}| \leq \sigma, \\ 0, & \text{otherwise.} \end{cases}$$

7 **Definition A.1.** *Noisy gradient.* Let f be differentiable and $\nabla_x f(x)$ denote the gradient of f
 8 in x . We consider that we only have access to $\nabla_x \hat{f}(x)$, a noisy estimator of $\nabla_x f(x)$ such that
 9 $\nabla_x \hat{f}(x) = \nabla_x f(x) + \varepsilon$ with $\varepsilon \in \mathbb{R}^{W \times H}$. We denote $\|\cdot\|_F$ as the Frobenius norm.

10 So far, we do not consider this noise to be random, and we aim at bounding the leftover noise after
 11 the application of our method. In particular, we'll notice that at the level σ^* , it only depends on the
 12 value of σ^* and is always upper bounded by the norm of the original noise.

13 **Proposition A.2.** *Let $\nabla \hat{f} = \nabla f + \varepsilon$ as the noisy gradient of f , with $\varepsilon \in \mathbb{R}^{W \times H}$. For $\sigma^* =$
 14 $\inf \{\sigma : \|\mathcal{F}(\nabla \hat{f}) \odot K^\sigma\|_F^2 = 0\}$, we have*

$$\|\mathcal{F}^{-1}(\mathcal{F}(\nabla \hat{f}) \odot K^{\sigma^*}) - \nabla f\|_F^2 = \|\mathcal{F}^{-1}(\mathcal{F}(\varepsilon) \odot K^{\sigma^*})\|_F^2 \leq \|\varepsilon\|_F^2, \quad (1)$$

15 where $\odot$ is the Hadamard product, K^{σ^*} a binary mask for low-pass filtering of frequency σ , and
 16 $\bar{K}^{\sigma^*}$ is the opposite mask.

Proof.

$$\|\mathcal{F}^{-1}(\mathcal{F}(\nabla \hat{\mathbf{f}} \odot \mathbf{K}^\sigma)) - \nabla \mathbf{f}\|_F^2 = \|\mathcal{F}^{-1}(\mathcal{F}(\nabla \hat{\mathbf{f}} \odot \mathbf{K}^\sigma)) - \mathcal{F}^{-1}(\mathcal{F}(\nabla \mathbf{f}))\|_F^2 \quad (2)$$

$$= \|\mathcal{F}^{-1}(\mathcal{F}(\nabla \mathbf{f} + \varepsilon \odot \mathbf{K}^\sigma)) - \mathcal{F}^{-1}(\mathcal{F}(\nabla \mathbf{f}))\|_F^2 \quad (3)$$

$$= \|\mathcal{F}^{-1}(\mathcal{F}(\nabla \mathbf{f} \odot \mathbf{K}^\sigma) + \mathcal{F}(\varepsilon \odot \mathbf{K}^\sigma)) - \mathcal{F}^{-1}(\mathcal{F}(\nabla \mathbf{f}))\|_F^2 \quad (4)$$

$$= \|\mathcal{F}^{-1}(\mathcal{F}(\nabla \mathbf{f} \odot \mathbf{K}^\sigma) + \mathcal{F}(\varepsilon \odot \mathbf{K}^\sigma) - \mathcal{F}(\nabla \mathbf{f}))\|_F^2 \quad (5)$$

$$= \|\mathcal{F}^{-1}(\mathcal{F}(\varepsilon \odot \mathbf{K}^\sigma) - \mathcal{F}(\nabla \mathbf{f} \odot \bar{\mathbf{K}}^\sigma))\|_F^2 \quad (6)$$

$$(7)$$

17 Then by selecting $\sigma^* = \arg \min_{\sigma} \|\mathcal{F}(\nabla \mathbf{f} \odot \bar{\mathbf{K}}^\sigma) - \mathcal{F}(\nabla \mathbf{f})\|_F^2 = 0$, we have that

$$\|\mathcal{F}^{-1}(\mathcal{F}(\varepsilon \odot \mathbf{K}^\sigma) - \mathcal{F}(\nabla \mathbf{f} \odot \bar{\mathbf{K}}^\sigma))\|_F^2 = \|\mathcal{F}^{-1}(\mathcal{F}(\varepsilon) \odot \mathbf{K}^\sigma)\|_F^2 \quad (8)$$

$$\leq \|\varepsilon\|_F^2 \quad (9)$$

$$(10)$$

18

□

19 Now, by adding an assumption of randomness on the noise ε , we can deduce the distribution of the
20 ratio of the two last members of the previous proposition, and deduce an order of scale of the norm of
21 the remaining noise after filtration, compared to its original norm.

22 **Proposition A.3.** *Let the noise $\varepsilon \in \mathbb{R}^{W \times H}$ follow a normal distribution $\varepsilon \sim \mathcal{N}(0, \varsigma)^{\otimes WH}$. Then
23 the norm of the Fourier spectra of the noise $\|\mathcal{F}(\varepsilon)\|_F^2 \sim \Gamma(k = 2WH, \theta = \varsigma^2 WH)$ and filtered
24 noise $\|\mathcal{F}(\varepsilon) \odot \mathbf{K}^\sigma\|_F^2 \sim \Gamma(k = 8\sigma^2, \theta = 4\varsigma^2 \sigma^2)$ follows Gamma distributions.*

25 Therefore, the ratio of the two distributions $R = \|\mathcal{F}(\varepsilon)\|_F^2 / \|\mathcal{F}(\varepsilon) \odot \mathbf{K}^\sigma\|_F^2$ follows a Beta prime
26 distribution $R \sim \beta'(2WH, 4\sigma^2, 1, \frac{WH}{4\sigma^2})$.

27 *Proof.* As defined previously, $\varepsilon_{(i,j)} \sim \mathcal{N}(0, \varsigma)$. It is well known that the Fourier transform of that
28 random matrix follows a complex normal distribution, or equivalently the real and imaginary random
29 variables of the Fourier transform are i.i.d normally distributed, as $\mathcal{RF}(\varepsilon)_{(i,j)} \sim \mathcal{N}(0, \varsigma\sqrt{WH/2})$
30 and $\mathcal{IF}(\varepsilon)_{(i,j)} \sim \mathcal{N}(0, \varsigma\sqrt{WH/2})$. Therefore the scaled norm follows a χ^2 distribution.

$$\frac{\|\mathcal{F}(\varepsilon)\|_F^2}{\varsigma\sqrt{WH/2}} \sim \chi_{2WH}^2 \quad \text{and equivalently,} \quad \frac{\|\mathcal{F}(\varepsilon) \odot \mathbf{K}^\sigma\|_F^2}{\varsigma\sqrt{WH/2}} \sim \chi_{8\sigma^2}^2. \quad (11)$$

31 Note that χ^2 distributions are a specific case of Gamma distributions, and can therefore be noted as
32 follows $\|\mathcal{F}(\varepsilon)\|_F^2 \sim \Gamma(k = WH, \theta = \varsigma^2 WH)$ and filtered noise $\|\mathcal{F}(\varepsilon) \odot \mathbf{K}^\sigma\|_F^2 \sim \Gamma(k = 4\sigma^2, \theta =$
33 $4\varsigma^2 \sigma^2)$. Finally, let $R = \|\mathcal{F}(\varepsilon)\|_F^2 / \|\mathcal{F}(\varepsilon) \odot \mathbf{K}^\sigma\|_F^2$. Since both random variables follow a Gamma
34 distribution, their ratio is a Beta prime distribution of parameters $R \sim \beta'(WH, 4\sigma^2, 1, \frac{WH}{4\sigma^2})$.

35 The results remain valid in the input space by Parseval's Theorem [1], up to a constant factor. □

36 **Proposition A.4.** *We recall that SmoothGrad is defined as $SG = \frac{1}{n} \sum_{i=1}^n \nabla_{\mathbf{x}} \mathbf{f}(\mathbf{x} + \delta_i)$ with
37 $\forall_{i=1, \dots, n} \delta_i \in \mathcal{N}(0, \varsigma)^{\otimes WH}$. $\nabla_{\mathbf{x}} \mathbf{f}(\mathbf{x} + \delta_i)$ is a random matrix and therefore SG is. Assuming our
38 predictor $\mathbf{f} \in L\text{-Lip}(\mathcal{X})$ is L -Lipschitz. We denote $\|\cdot\|_2$ as the spectral norm, and define the variance
39 as $\mathbb{V}(SG) = \max(\|\mathbb{E}((SG - \mathbb{E}SG) \cdot (SG - \mathbb{E}SG)^T)\|_2, \|\mathbb{E}((SG - \mathbb{E}SG)^T \cdot (SG - \mathbb{E}SG))\|_2)$.
40 We then have, for $t > 0$,*

$$\mathbb{P}(\|SG - \mathbb{E}SG\|_2 \geq t) \leq (W + H) \cdot \exp\left(\frac{-t^2 n^2 / 2}{\mathbb{V}(SG) + 2Lt/3}\right). \quad (12)$$

41 *Proof.* We first demonstrate the bounded difference property

$$\forall_{i=1 \dots n} \|\nabla_{\mathbf{x}} \mathbf{f}(\mathbf{x} + \delta_i) - \mathbb{E} \nabla_{\mathbf{x}} \mathbf{f}(\mathbf{x} + \delta_i)\|_2 \leq \|\nabla_{\mathbf{x}} \mathbf{f}(\mathbf{x} + \delta_i)\|_2 + \|\mathbb{E} \nabla_{\mathbf{x}} \mathbf{f}(\mathbf{x} + \delta_i)\|_2 \leq 2L. \quad (13)$$

42 Finally, the result follows by application of the Matrix Bernstein Inequality [2]. □

43 B Different measures of complexity

44 To corroborate the results in S3.1 and S3.3 computing the high-frequency content from Kolmogorov
45 image compression metric, we reproduce those experiments using the Laplace2-operator. We
46 additionally show the evolution of high-frequency content in VGG16 and ViT using both metrics. On
47 all the plots, we observe a similar trend between the metrics, confirming that :

- 48 • Gradient-based and prediction-based methods have a different signature in the Fourier
49 spectrum (see Figure 1)
- 50 • On both convolutional models (ResNet50 and VGG) we observe more high frequencies
51 when the models contain an MaxPooling operation or strided convolution instead of Aver-
52 agePooling (see Figures 2, 3, 3, top curves)
- 53 • On both convolutional models, training the model doesn't alleviate the high frequency
54 content (see Figures 2, 3, 3, bottom curves). However, we observe a different behaviour on
55 ViT, suggesting that a different mechanism is responsible for the high-frequency content in
56 this kind of architectures (see Figures 5, 6).

57 B.1 For the methods

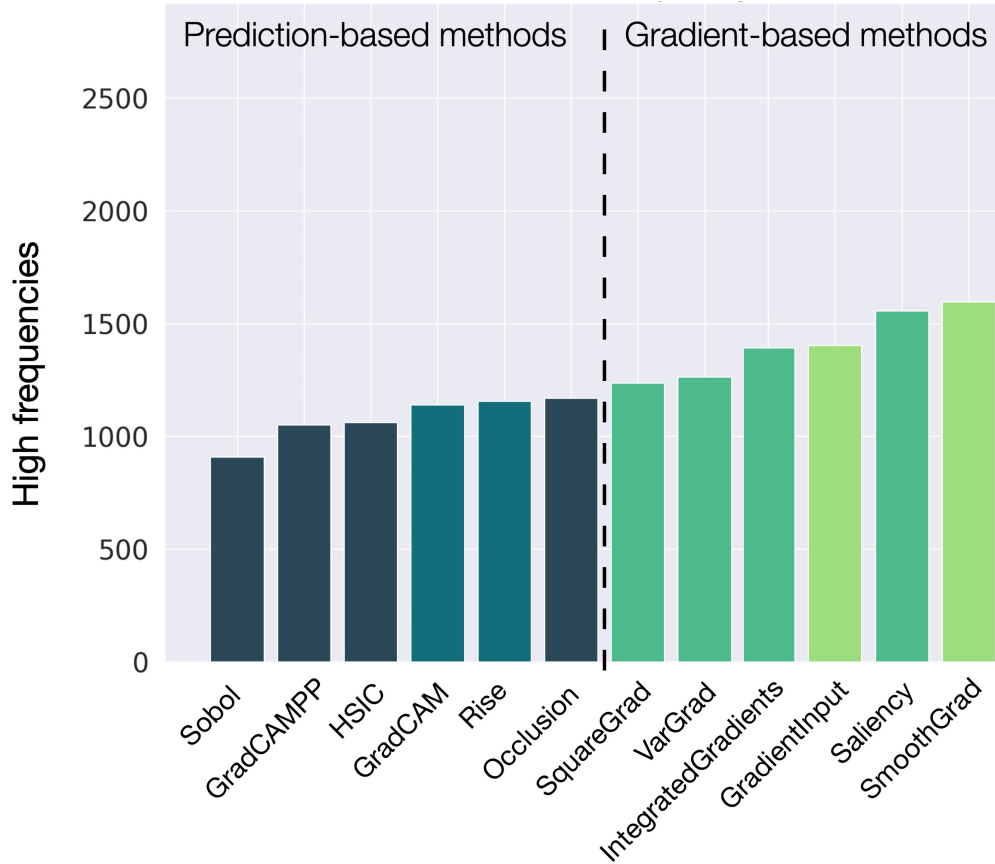


Figure 1: **High-frequency power in attribution methods.** High-frequency power present in the importance maps derived from different attribution methods, computed using Laplacian2 operator. Prediction-based methods produce less high-frequency content than gradient-based methods.

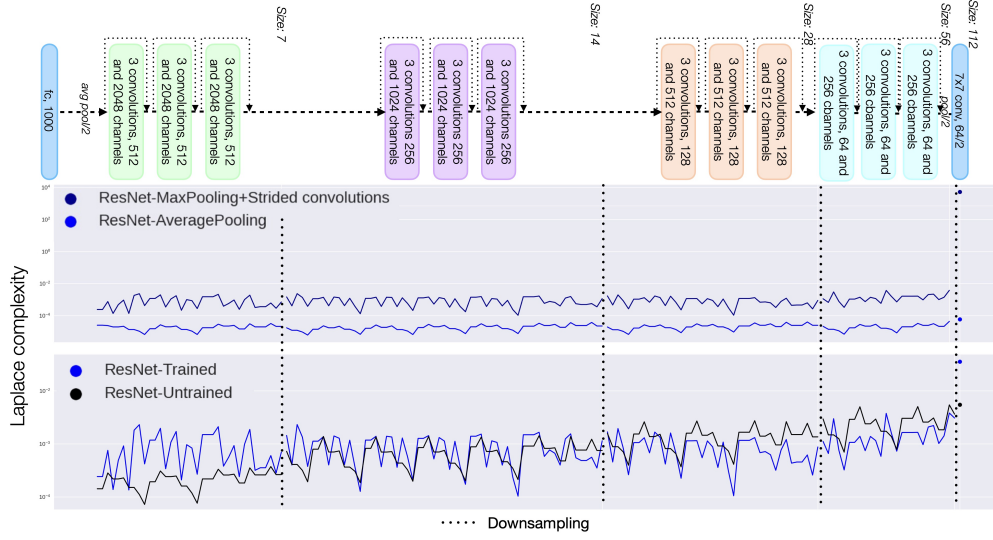


Figure 2: **Evolution of the high-frequency content in Resnet50.** We compute the high-frequency content along the depth of a ResNet50 varying either the weights or the pooling, using Laplace2-operator. The top curve illustrates the impact of different poolings, with MaxPooling and stride shown in dark blue and AveragePooling in light blue. The bottom curve represents the trained model, indicated by the blue curve, while the untrained model is represented by the black curve. Each point on the graph corresponds to a layer within the model.

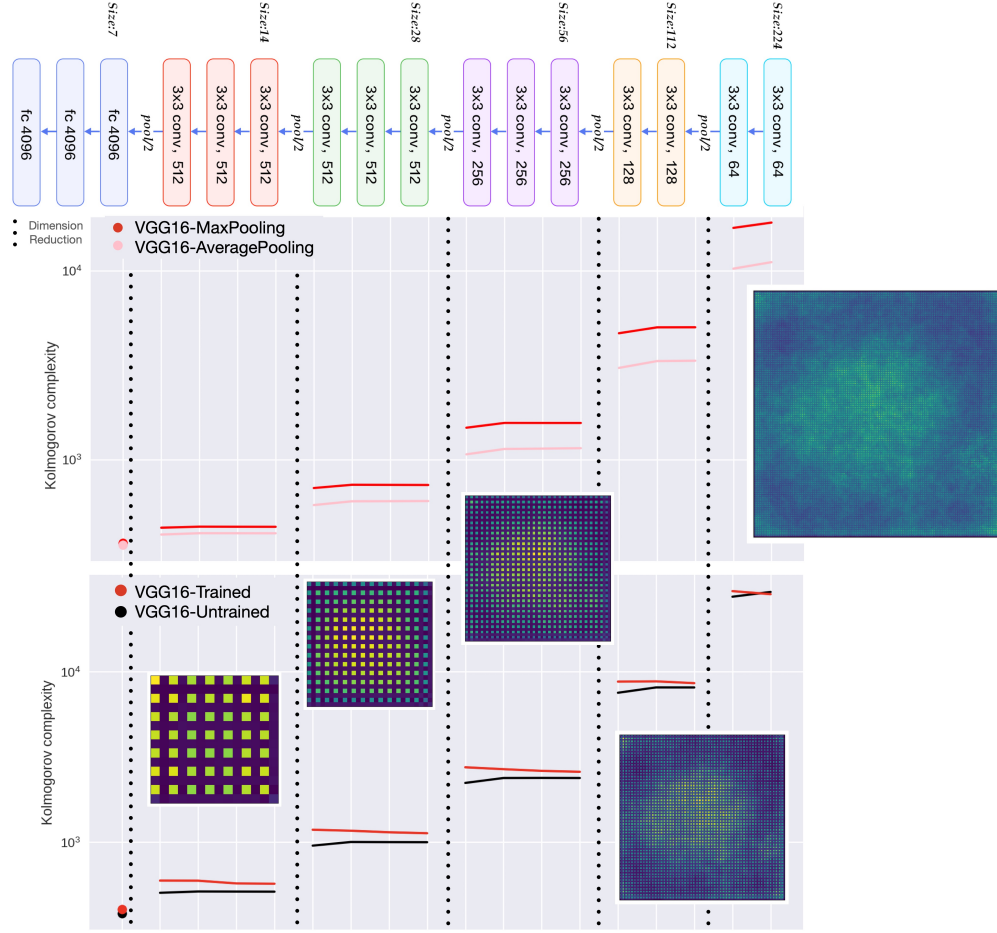


Figure 3: **Evolution of the high-frequency content in VGG16 with Kolmogorov complexity.** We compute the high-frequency content along the depth of a VGG16 varying either the weights or the pooling, using Kolmogorov image compression. The top curve illustrates the impact of different poolings, with MaxPooling shown in dark red and AveragePooling in pink. The bottom curve represents the trained model, indicated by the red curve, while the untrained model is represented by the black curve. Each point on the graph corresponds to a layer within the model.

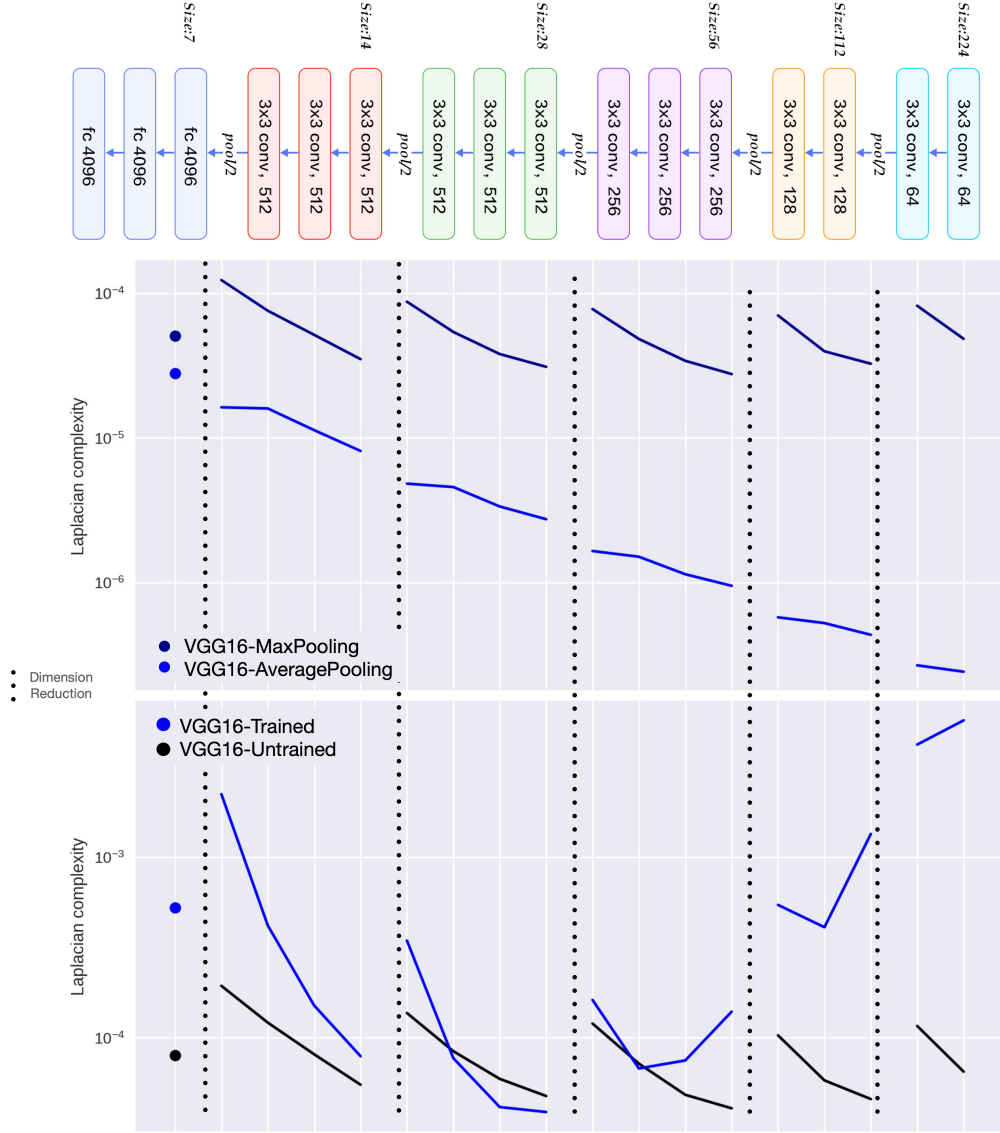


Figure 4: **Evolution of the high-frequency content in VGG16 with Laplace complexity.** We compute the high-frequency content along the depth of a VGG16 varying either the weights or the pooling, using Laplace2-operator. The top curve illustrates the impact of different poolings, with MaxPooling shown in dark blue and AveragePooling in light blue. The bottom curve represents the trained model, indicated by the blue curve, while the untrained model is represented by the black curve. Each point on the graph corresponds to a layer within the model.

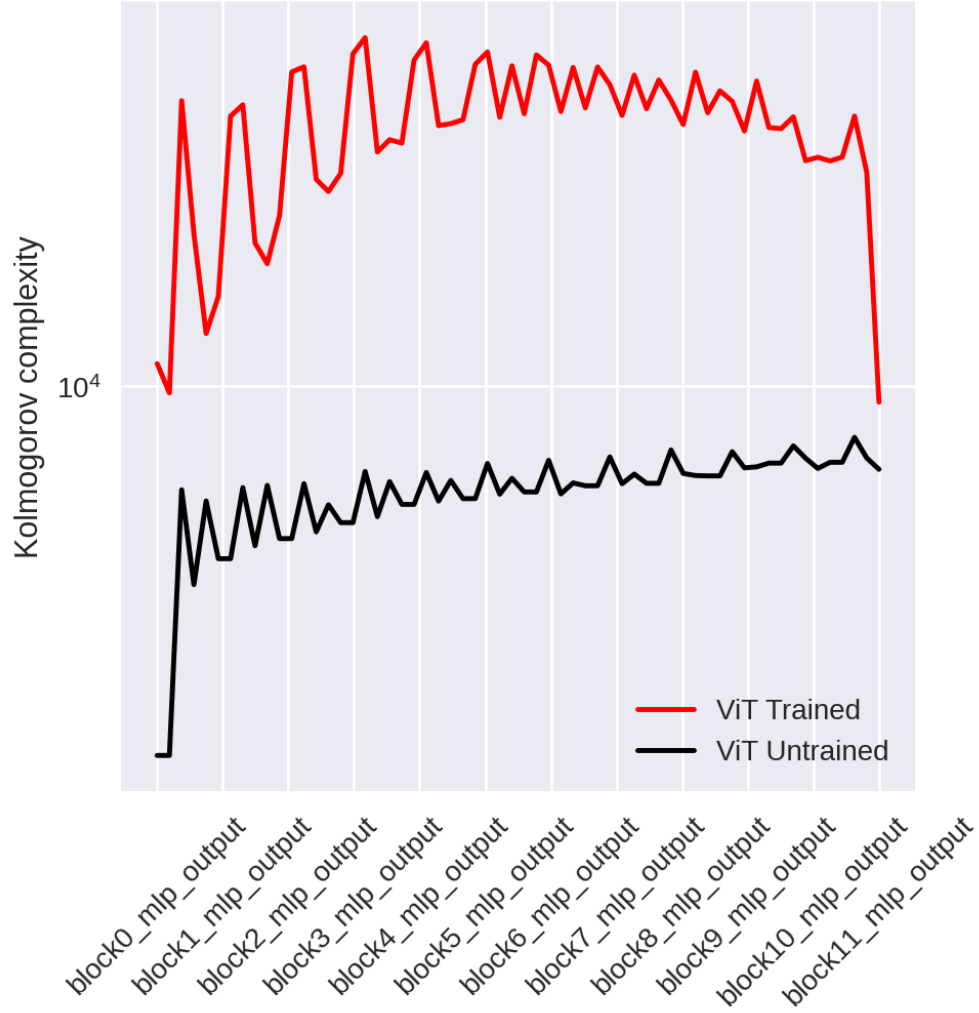


Figure 5: **Evolution of the high-frequency content in ViT with Kolmogorov complexity.** We compute the high-frequency content along the depth of a ViT varying the weights, using Kolmogorov image compression. The curve represents the trained model, indicated by the red curve, while the untrained model is represented by the black curve. Each point on the graph corresponds to a layer within the model. We show the labels representing the end of a block to give an idea of the evolution of the complexity inside a block.

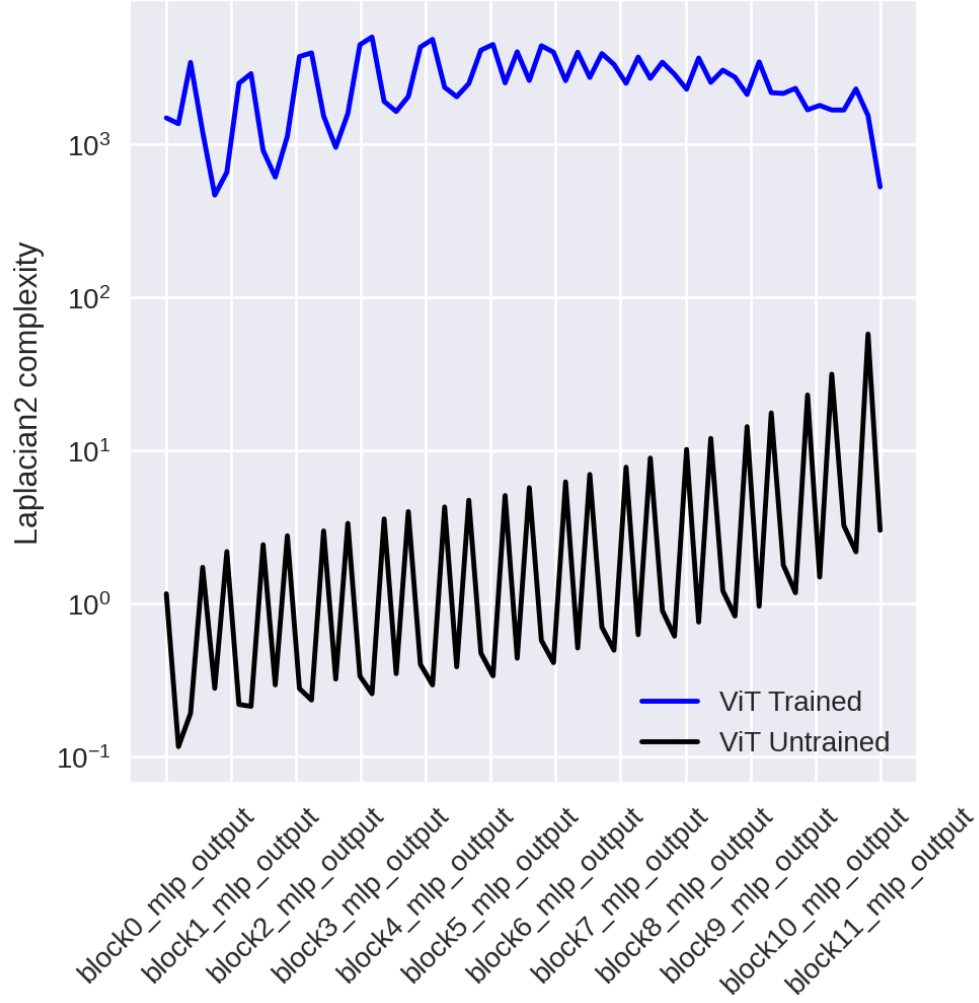


Figure 6: **Evolution of the high-frequency content in ViT with Laplace complexity.** We compute the high-frequency content along the depth of a ViT varying the weights, using Laplace2-operator. The curve represents the trained model, indicated by the blue curve, while the untrained model is represented by the black curve. Each point on the graph corresponds to a layer within the model. We show the labels representing the end of a block to give an idea of the evolution of the complexity inside a block.

59 C Control conditions for the evidence of noise in the gradient

60 We add several control conditions on the experiment S3.2 to show the difference between an informa-
 61 tive and uninformative gradient. We propose three different approaches :

- 62 • $\sigma = 0 \equiv \nabla_x^\sigma \mathbf{f}(\mathbf{x}) = 0$. In that case, we measure $\|\mathbf{f}(\mathbf{x} + \epsilon) - \mathbf{f}(\mathbf{x})\|_2$
- 63 • $\nabla_x^C \mathbf{f}(\mathbf{x}) = \rho(\nabla_x \mathbf{f}(\mathbf{x}))$ where ρ represents the permutation operator. Here we destroy the
 64 spatial structure, and therefore the information represented by high and low frequencies.
- 65 • $\nabla_x^C \mathbf{f}(\mathbf{x}) \sim \mathcal{U}(\min(\nabla_x \mathbf{f}(\mathbf{x})), \max(\nabla_x \mathbf{f}(\mathbf{x})))$. Finally, we measure the information
 66 carried by some random noise, following a uniform distribution.

67 In the three cases, we compute the L2 norm between the first-order approximation of the model and
 68 something else containing no relevant information. We therefore expect a resulting curve with a
 69 higher estimation error than the others, containing some relevant information. We observe that is
 70 always the case for ResNet and ViT. This observation is a bit less clear for some radius value in the
 two first conditions on ConvNeXT but is present in the third condition.

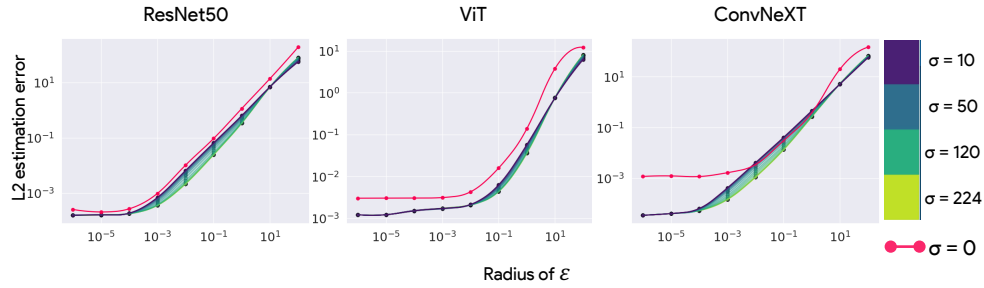


Figure 7: **Control condition** $\sigma = 0$. We plot the residual of the first-order approximation of the model, that is $\mathbf{f}(\mathbf{x} + \epsilon) \approx \mathbf{f}(\mathbf{x}) + \epsilon \nabla_x \mathbf{f}(\mathbf{x})$, with the gradient $\nabla \mathbf{f}$ filtered at different bandwidths σ . Additionally, we plot the control condition $\sigma = 0$ in pink, representing the absence of information from the gradient.

71

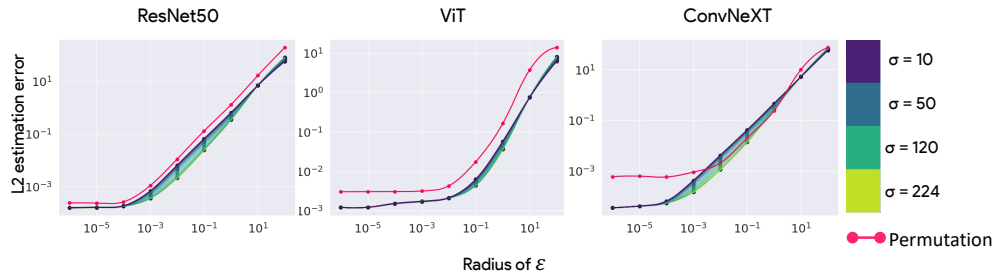


Figure 8: **Control condition** $\nabla_x^C \mathbf{f}(\mathbf{x}) = \rho(\nabla_x \mathbf{f}(\mathbf{x}))$. We plot the residual of the first-order approximation of the model, that is $\mathbf{f}(\mathbf{x} + \epsilon) \approx \mathbf{f}(\mathbf{x}) + \epsilon \nabla_x \mathbf{f}(\mathbf{x})$, with the gradient $\nabla \mathbf{f}$ filtered at different bandwidths σ . Additionally, we plot the control condition $\nabla_x^C \mathbf{f}(\mathbf{x}) = \rho(\nabla_x \mathbf{f}(\mathbf{x}))$ in pink, representing some unstructured information from the gradient.

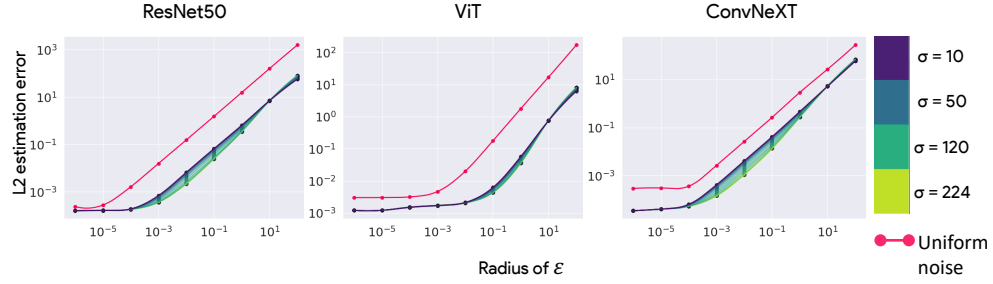


Figure 9: **Control condition** $\nabla_x^C \mathbf{f}(\mathbf{x}) \sim \mathcal{U}(\min(\nabla_x \mathbf{f}(\mathbf{x})), \max(\nabla_x \mathbf{f}(\mathbf{x})))$. We plot the residual of the first-order approximation of the model, that is $\mathbf{f}(\mathbf{x} + \varepsilon) \approx \mathbf{f}(\mathbf{x}) + \varepsilon \nabla_x \mathbf{f}(\mathbf{x})$, with the gradient $\nabla \mathbf{f}$ filtered at different bandwidths σ . Additionally, we plot the control condition $\nabla_x^C \mathbf{f}(\mathbf{x}) \sim \mathcal{U}(\min(\nabla_x \mathbf{f}(\mathbf{x})), \max(\nabla_x \mathbf{f}(\mathbf{x})))$ in pink, representing some random information from the gradient.

72 **References**

- 73 [1] Marc-Antoine Parseval. Mémoire sur les séries et sur l'intégration complète d'une équation aux différences
74 partielles linéaires du second ordre, à coefficients constants. *Mém. prés. par divers savants, Acad. des*
75 *Sciences, Paris,(1)*, 1:638–648, 1806.
- 76 [2] Joel A Tropp et al. An introduction to matrix concentration inequalities. *Foundations and Trends® in*
77 *Machine Learning*, 8(1-2):1–230, 2015.