

Eva Optimizer

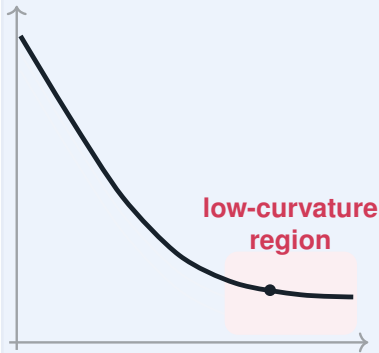
Escaping Low-Curvature Traps in Deep Learning

Preserve Adam if its step matches the local gradient scale;
correct only mismatched coordinates.

A. Di Cecco · C. Metta · A. Papini · M. Fantozzi · S. G. Galfrè ·
M. Vegliò · **L. A. Bianchi** · M. Parton · F. Morandin

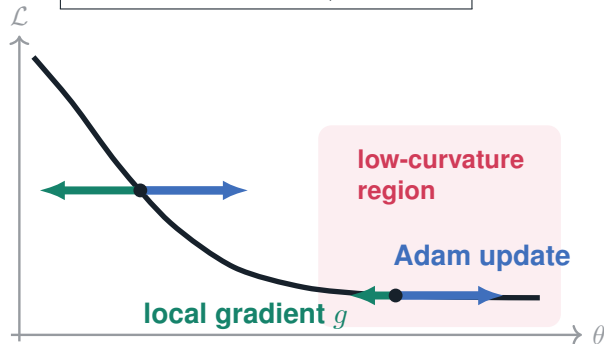
UCBM, CNR, Chalmers, Parma, Pisa, **Trento**, Chieti-Pescara; **CuriosAI**

2026.08.18, **ICPR 2026**



$$\Delta\theta_{\text{Adam},i}^{(t)} = -\alpha \frac{\widehat{m}_i^{(t)}}{\sqrt{\widehat{v}_i^{(t)} + \varepsilon}}$$

$\widehat{m}_i^{(t)}$: moving average of recent gradients
 $\widehat{v}_i^{(t)}$: moving average of squared gradients



In a slowly varying flat direction:

$|g_{t-1,i}|$ can become small,

$|\Delta\theta_{\text{Adam},i}^{(t)}|$ need not shrink at the same rate.

One-dimensional intuition: low curvature is considered together with gradient stability.

Low-curvature / stable-gradient approximation

$$\begin{aligned}\widehat{m}_i^{(t)} &\simeq g_{t-1,i} \\ \sqrt{v_i^{(t)}} &\simeq |g_{t-1,i}| \\ \Rightarrow |\Delta\theta_{\text{Adam},i}^{(t)}| &\simeq \alpha\end{aligned}$$

$$\kappa_i^{(t)} = \frac{|\Delta\theta_{\text{Adam},i}^{(t)}|}{|g_{t-1,i}|}$$

Component-wise scale mismatch indicator associated with the low-curvature regime.

$\kappa_i \leq 1$ **KEEP ADAM**

$\kappa_i > 1$ **CORRECT**

Operationally: compare the two magnitudes directly—no division is required.

$$\Delta\theta_{\text{Eva},i}^{(t)} = \begin{cases} \Delta\theta_{\text{Adam},i}^{(t)}, & \kappa_i^{(t)} \leq 1, \\ \sigma \lambda g_{t-1,i}, & \kappa_i^{(t)} > 1 \end{cases}$$

Different coordinates, same iteration t



Only flagged coordinates change.

Adam is preserved everywhere else.

$\lambda = 0.1$
fixed default used in
reported experiments

Sign convention: $\theta_t = \theta_{t-1} + \Delta\theta_t$

First-order contribution: $g_{t,i}\Delta\theta_i = \sigma\lambda g_{t,i}g_{t-1,i}$

If $g_{t,i}g_{t-1,i} > 0$, the current and previous gradients are locally aligned.

Eva⁺: $\sigma = +1$

$$\Delta\theta_i = +\lambda g_{t-1,i}$$



intended role

leave a flagged region

Eva⁻: $\sigma = -1$

$$\Delta\theta_i = -\lambda g_{t-1,i}$$



intended role

reduce the flagged step

Separate variants: sign chosen before training; the gate tests magnitudes, not alignment.

When the gate triggers and $0 < \lambda \leq 1$:

$$|\Delta\theta_{\text{Eva},i}| = \lambda |g_{t-1,i}| \leq |g_{t-1,i}| < |\Delta\theta_{\text{Adam},i}|$$

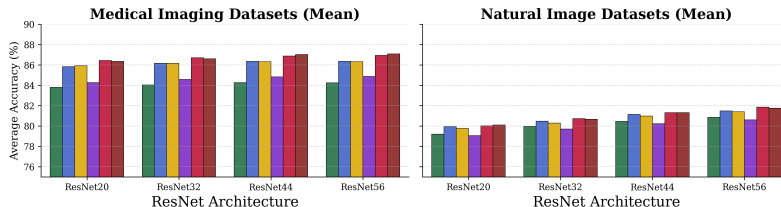
component-wise

no Hessian

element-wise $O(d)$

Implementation note: Eva needs access to g_{t-1} , typically via a previous-gradient buffer.

UNDER MATCHED HYPERPARAMETERS



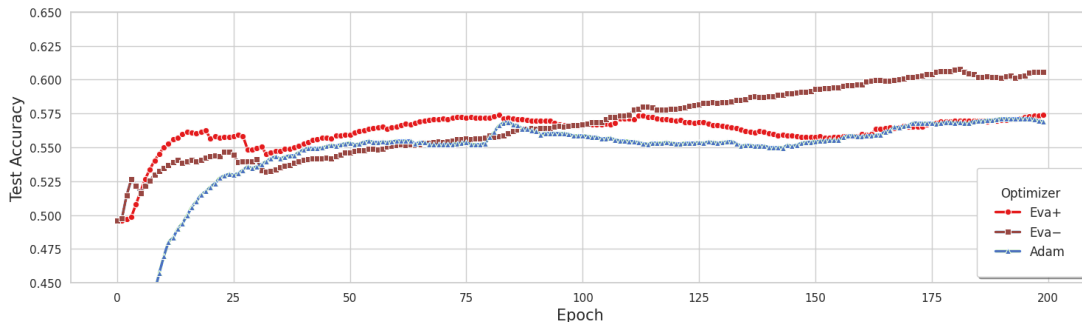
63/64
best or tied by
an Eva variant

SGD Adam AdamW RMSProp Eva⁺ Eva⁻
16 datasets × 4 ResNets · 5 independent runs

Reported sweep. Note that optimizer-specific tuning can alter rankings.

42 Eva⁺ unique best
17 Eva⁻ unique best
4 ties between Evas
1 Adam unique best

Representative run · RetinaMNIST · ResNet-56



Eva⁺
faster early progress

Eva⁻
accuracy in the long run

Pretrained vision

4 families × 2 datasets

ConvNeXt · EfficientNet · InceptionNet
MobileNetV2 on two medical datasets

VAE reconstruction

4/4

Eva⁻ best reported SSIM

Tabular classification

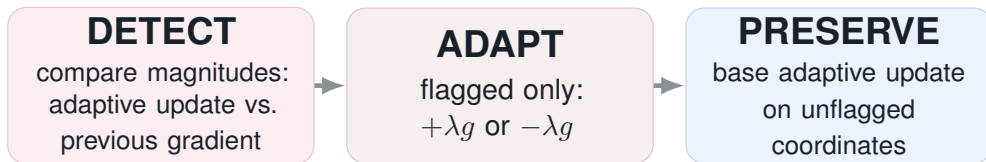
4/4

Eva⁺ best reported accuracy

BERT / IMDB

92.3% Eva⁺

vs AdamW 91.7% and Adam 91.5%



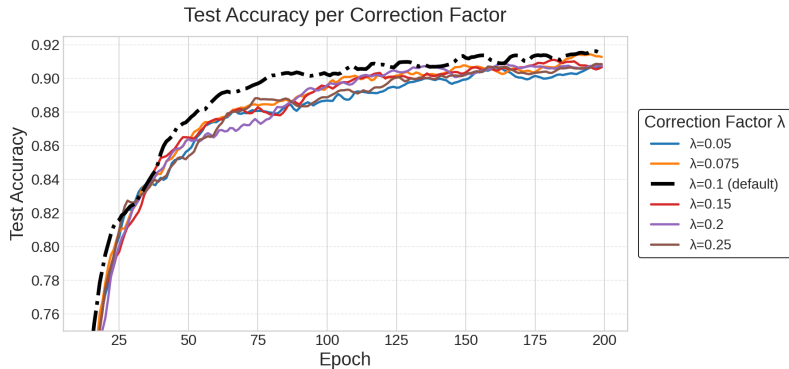
Eva changes only what the gate flags.

Flagged coordinates receive a controlled correction; unflagged coordinates use the Adam update.



CuriousAI/Eva on
GitHub

We do not need to trust Adam equally in every coordinate, at every step.

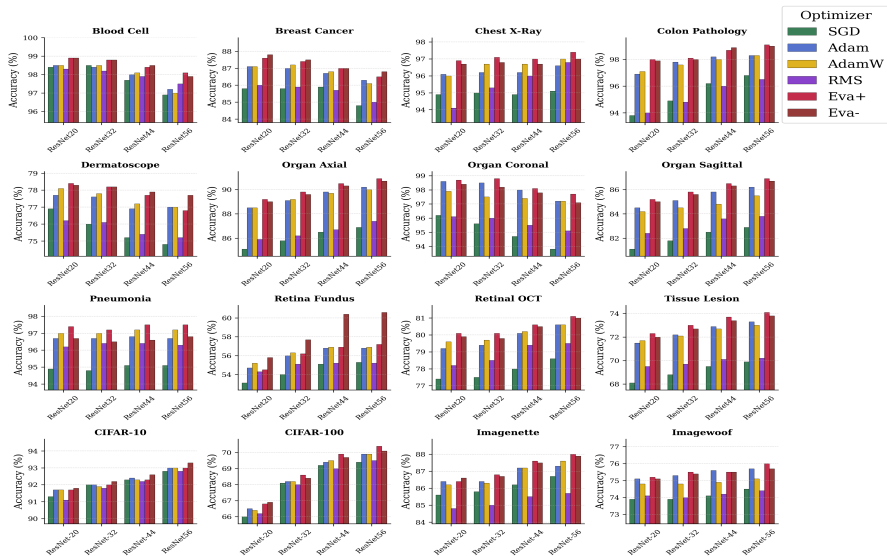


$$\lambda = 0.1$$

fixed default in the reported experiments.

This ResNet-20 / CIFAR-10 ablation supports the choice in this configuration, it is no proof of a universal optimum.

Detailed Optimizer Performance per Dataset



Assumptions along coordinate i

- low curvature;
- gradient stability: $g_{t,i} \approx g_{t-1,i}$;
- near-zero ideal update: $\Delta\theta_i^* \approx 0$.

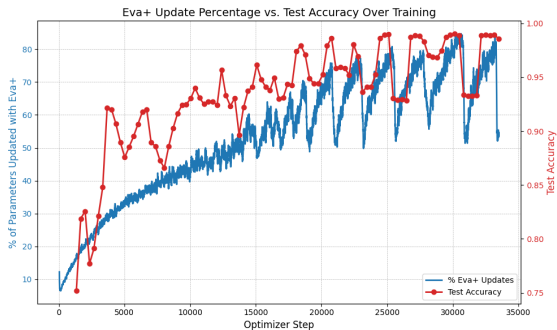
Local simplifying assumptions, not global training conditions.

$$|\Delta\theta_{\text{Adam},i}| \simeq \alpha$$

$$\text{gate active} \Rightarrow |g_{t-1,i}| < \alpha$$

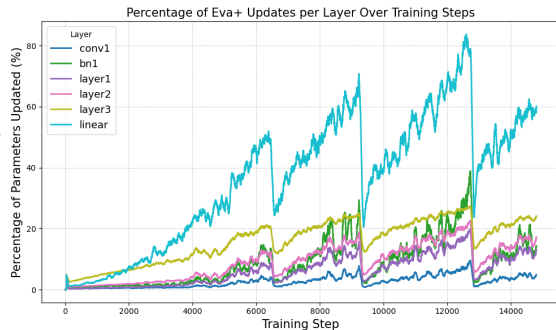
$$e_{\text{Eva}}^2 \approx \lambda^2 |g_{t-1,i}|^2 < \\ < \lambda^2 \alpha^2 < \alpha^2 \approx e_{\text{Adam}}^2$$

For $\lambda = 0.1$, this model puts the Eva squared-error scale below 1% of α^2 .

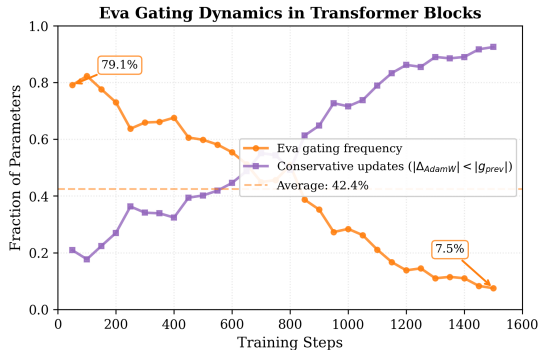
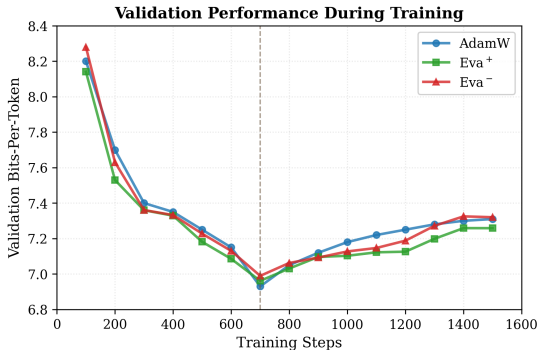


update usage vs. accuracy

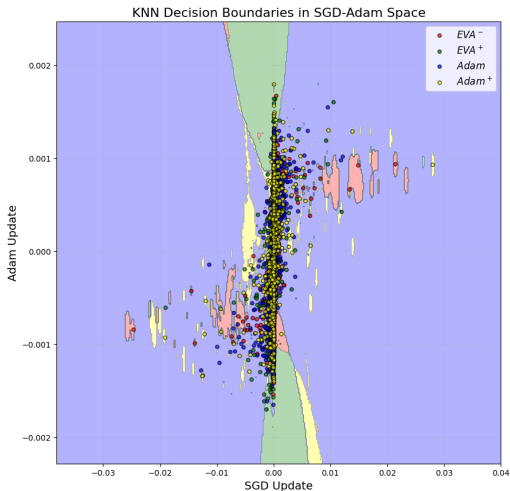
In the reported PathMNIST experiment, activation evolves through training and differs substantially by layer.



update usage across layers



Limitations: ~ 10 M-parameter Shakespeare model; selective gating; 50–50 blended update; heuristic blending ratio; no systematic ablation.



Protocol

50,000 recorded parameter-level triplets

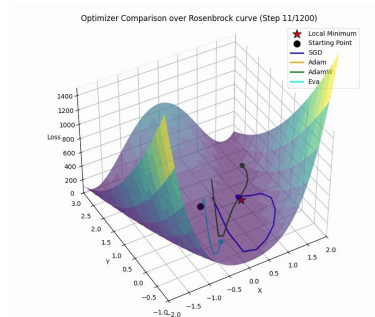
Candidate update labels: Adam, reversed Adam, SGD, reversed SGD

k -NN visualization with $k = 500$

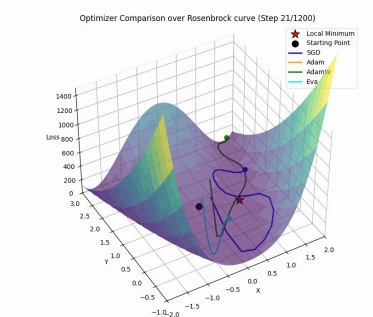
Reported pattern

Eva^+ -like choices appear where Adam updates are large relative to gradient/SGD updates; Adam dominates where the scales are proportional.

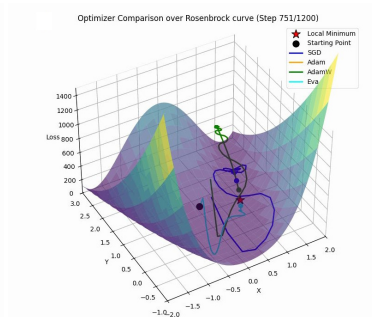
Only a supportive empirical analysis, not a proof.



step 11



step 21



step 751

Toy illustration: controlled Eva trajectories versus large Adam/AdamW excursions on the Rosenbrock surface.