

Table R1: (Reviewer m6P2: Weakness 1,2, Reviewer 41nW: Weakness 1,2) Comparison of Spearman’s ρ between the robust accuracies and the proxy values on CIFAR-10 in NAS-Bench-201 search space. All robust accuracies are obtained against diverse adversarial attacks, where 300 models are randomly selected in the NAS-Bench-201 search space and adversarially trained with PGD-7 attack. Avg. stands for average Spearman’s ρ values with all accuracies within each task.

Proxy Type	FGSM	PGD	CW	DeepFool	SPSA	LGV	AutoAttack	Avg.
FLOPs	0.357	0.446	0.189	0.364	0.196	0.347	0.365	0.323
#Params.	0.367	0.457	0.182	0.371	0.209	0.355	0.375	0.331
Plain	-0.006	-0.028	0.072	0.009	0.009	0.010	0.015	0.012
Grasp	0.397	0.420	0.179	0.393	0.249	0.401	0.397	0.348
Fisher	0.233	0.279	0.234	0.242	0.092	0.239	0.244	0.223
GradNorm	0.378	0.446	0.264	0.421	0.149	0.401	0.405	0.352
SynFlow	0.369	0.442	0.202	0.397	0.196	0.387	0.383	0.339
NASWOT	0.311	0.354	0.240	0.250	0.197	0.265	0.280	0.271
CRoZe	0.441	0.532	0.220	0.454	0.240	0.449	0.458	0.399

Table R2: (Reviewer m6P2: Weakness 2, Reviewer DQhw: Weakness 2, Reviewer 41nW: Weakness 3) Comparisons of the final performance of the searched network in DARTS search space on CIFAR-10. All models are adversarially trained with PGD-7 attack.

NAS Method	NAS Type	Search Cost (GPU sec)	Clean	Robustness					
				PGD-20	CW	SPSA	LGV	AutoAttack	Avg.
DrNAS	Clean one-shot	46857	86.45	54.66	6.44	85.86	79.64	52.40	55.60
AdvRush	Robust one-shot	251245	85.98	53.89	6.68	84.81	79.61	51.88	55.57
GradNorm SynFlow	Clean zero-shot	9740	81.61	49.86	12.02	77.19	73.27	46.69	51.61
		10138	77.08	45.95	26.50	75.78	74.14	42.45	52.96
CRoZe	Robust zero-shot	17066	85.05	52.04	16.82	83.23	77.62	49.15	55.77

Fig. R1: End-to-end average robust accuracy on CIFAR-10.

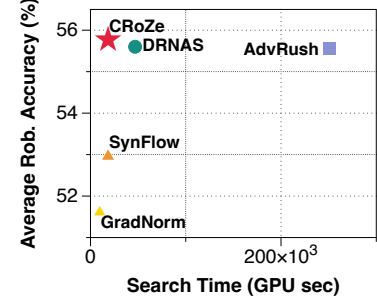


Table R3: (Reviewer DQhw: Weakness 1) Comparison of Spearman’s ρ between the final accuracies and the proxy values on the NAS-Bench-201 search space with various weight initialization methods. All models are trained with both standard and adversarial training on CIFAR-10.

Weight Initialization Type	Standard-trained					Adversarially-trained							
	Clean	FGSM	PGD	CC.	Avg.	FGSM	PGD	CW	DeepFool	SPSA	LGV	AutoAttack	Avg.
Random	0.823	0.826	0.188	0.436	0.568	0.441	0.532	0.220	0.454	0.240	0.449	0.458	0.399
Kaiming	0.812	0.818	0.189	0.430	0.562	0.428	0.512	0.217	0.443	0.227	0.426	0.436	0.384
Xavier	0.816	0.822	0.190	0.433	0.565	0.428	0.513	0.217	0.442	0.227	0.425	0.436	0.384

Table R4: (Reviewer 41nW: Weakness 3) Comparisons of the final performance between the searched network from our proxy and human-made architectures on CIFAR-10. All models are adversarially trained with PGD-7 attack.

Model	#Params (MB)	Clean	PGD-20	HRS
Conv4	0.03	60.12	29.98	40.01
Conv6	0.05	65.92	33.04	44.02
MobileNetV2	2.24	66.04	33.04	46.79
ResNet12	8.00	82.94	49.69	62.15
CRoZe	5.52	85.05	52.04	64.57

Table R5: (Reviewer KEHv: Weakness 2) Comparisons of the final performance of the searched network in NAS-Bench-201 search space on CIFAR-10 and CIFAR-100 with diverse existing zero-cost proxies.

Proxy Type	CIFAR-10				CIFAR-100			
	Clean	FGSM	PGD	CC.	Clean	FGSM	PGD	CC.
NASWOT	92.96	59.90	41.70	35.19	70.03	17.00	4.40	16.12
NASI (T)	93.08	62.60	41.10	34.99	69.51	17.80	3.70	16.77
NASI (4T)	93.55	64.90	44.00	36.12	71.20	23.10	8.10	19.73
Eigen-NAS (k=20)	93.46	59.60	36.80	36.75	71.42	24.30	10.00	20.42
KNAS	93.38	63.80	44.90	34.54	70.78	21.50	6.60	18.81
CRoZe	93.70	68.00	48.20	38.83	71.80	27.00	10.90	20.09

Table R6: (Reviewer m6P2: Question 3) Comparison of Spearman’s ρ between the final accuracies and the proxy values on CIFAR-10 in NAS-Bench-201 search space. + means the usage of ensembling zero-cost proxies with ours.

Proxy Type	CIFAR-10				CIFAR-100				ImageNet16-120		
	Clean	FGSM	PGD	CC.	Clean	FGSM	PGD	CC.	Clean	FGSM	PGD
CRoZe	0.823	0.826	0.188	0.823	0.784	0.786	0.343	0.784	0.765	0.596	0.707
Ensemble	0.803	0.681	0.543	0.771	0.793	0.739	0.444	0.798	0.749	0.638	0.296
CRoZe + Ensemble	0.894	0.872	0.633	0.894	0.894	0.851	0.415	0.878	0.810	0.688	0.259