

Semantic and Instance Segmentation

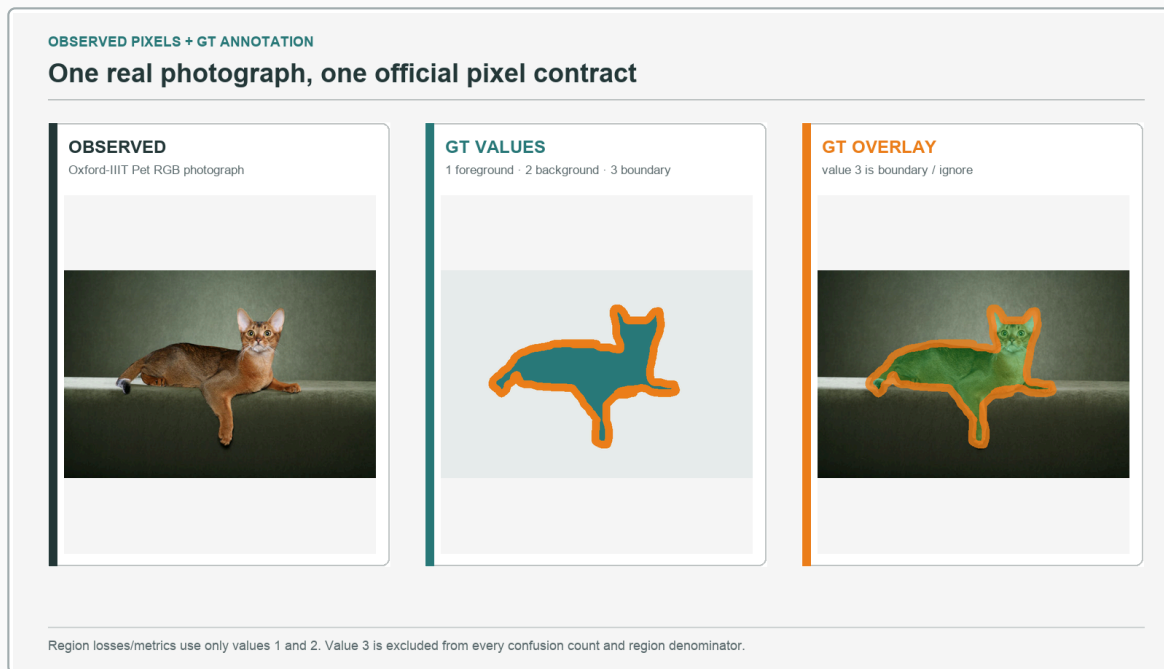
Pixel evidence becomes a class field—and, when needed, separate object masks

Prof. Nipun Batra

ES 667 · Deep Learning · IIT Gandhinagar

**Detection drew rectangles;
segmentation must assign
pixels**

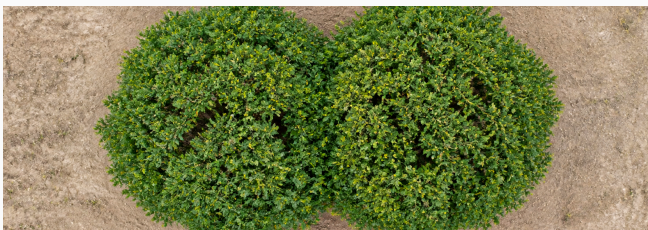
One real photograph changes the supervision from a box to pixels



OBSERVED Oxford pixels + official GT trimap. Value 1 is foreground, 2 background, and 3 boundary/ignore; no model output appears here.

Commit: which output can answer both requirements?

QUESTION



HELD CASE · GENERATED CONCEPT SCENE Two touching crowns form one region. Report total canopy area **and** count two trees.

A

one global class

B

one class + one
box

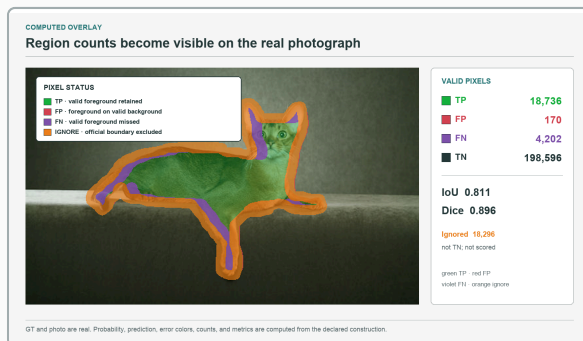
C

one $H \times W$
foreground map

D

a scored set of
separate object
masks

Commit silently; the same scene returns after semantic and instance outputs.



OBSERVED pixels + GT ·
CONSTRUCTED probabilities ·
COMPUTED errors

01 · OUTPUT CONTRACT

Separate semantic labels, instances, and panoptic records.

02 · PIXEL LEARNING

Turn logits into probabilities; train with declared losses.

03 · SPATIAL DETAIL

Decode resolution with skips, upsampling, and alignment.

04 · METRICS

Threshold once; compute IoU, Dice, and boundary failures.

05 · INSTANCES

Predict a scored set of separate object masks.

06 · EVALUATION

Write void, empty-mask, matching, and averaging policies.

— truth / TRAIN — logit, probability / INFER — hard prediction — overlap / TP — FP or FN — EVAL

Three clocks prevent a probability from becoming a metric

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

TRAIN compares with labels; INFER makes an output; EVAL scores that output under a declared protocol.

The exact arithmetic case is constructed—not an orchard annotation

V VISUAL

generated concept



not a measured orchard sample

constructed truth G

0	1	1	0
0	1	1	0
0	1	1	0
0	0	0	0

$|G| = 6$ pixels

constructed identities

0	1	2	0
0	1	2	0
0	1	2	0
0	0	0	0

IDs withheld from G .

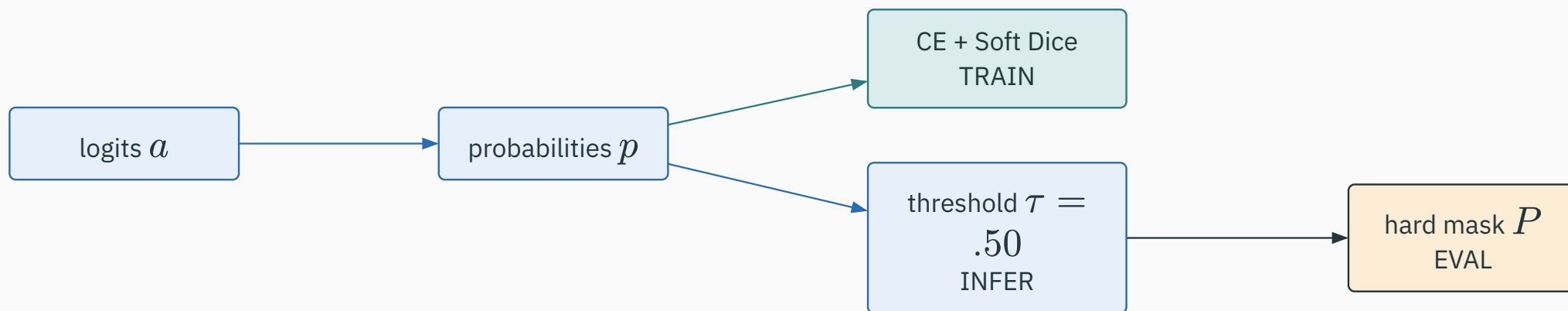
The 4×4 ledger isolates the arithmetic: one union supports area; two identities support counting.

Semantic, instance, and panoptic outputs answer different questions

Task	Output contract	Question answered
semantic	$H \times W$ class labels	Which class owns this pixel?
instance	variable scored set of object masks	Which pixels belong to object j ?
panoptic	$H \times W$ class–identity pairs	What and which object, everywhere?

“Stuff” such as sky usually needs a class; countable “things” such as trees also need an identity.

One prediction will cross all three clocks



HELD CASE · PROMISE The exact G , p , threshold, and P will remain fixed through loss, inference, and evaluation.

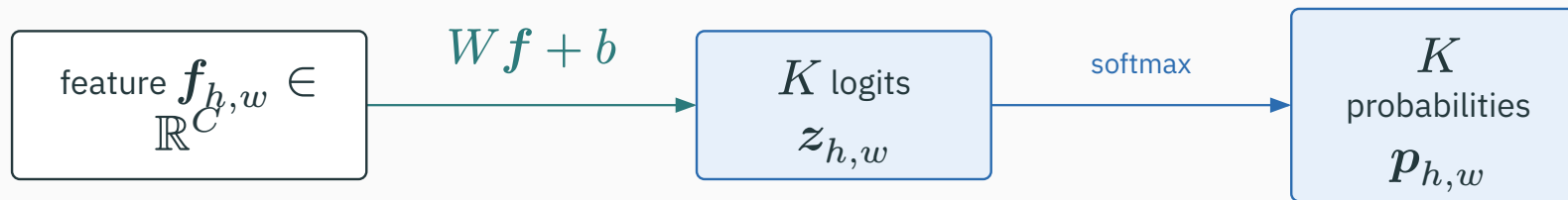
**TRAIN: classify every spatial
location**

A semantic head emits one score vector per pixel

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP



Segmentation shares one classifier across locations; evidence changes, weights do not. Argmax waits for INFER.

Softmax acts across classes at one location—not across pixels

D DERIVATION

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

HELD CASE • ONE PIXEL Classes are **background, road, car** and $\mathbf{z} = \begin{pmatrix} .2 \\ 1.4 \\ -.3 \end{pmatrix}$.

EXPONENTIATE

NORMALIZE

$$\exp(\mathbf{z}) \approx \begin{pmatrix} 1.22 \\ 4.05 \\ .74 \end{pmatrix} \quad \sum_k \exp(z_k) \approx 6.01$$

$$\mathbf{p} \approx \begin{pmatrix} .203 \\ .674 \\ .123 \end{pmatrix} \quad p_{\text{road}} = .674$$

$$\ell = -\log p_{\text{road}} = -\log .674 \approx .395.$$

$$\sum_k p_{h,w,k} = 1 \text{ independently at every } (h, w).$$

Binary foreground can use one logit and one sigmoid

D DERIVATION

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

$$p_i = \sigma(a_i) = \frac{1}{1+e^{-a_i}}, \quad a_i = \log \frac{p_i}{1-p_i}.$$

p	.90	.60	.30	.10	.05
$a =$ $\text{logit}(p)$	2.197	.405	−.847	−2.197	−2.944

The held binary mask uses one foreground probability. A K -class semantic task instead keeps K logits and softmax.

The held foreground probabilities are now fixed

TRAIN probability $\leftrightarrow$ truth; CE / Dice**INFER** probability $\rightarrow$ label / object masks**EVAL** hard output $\leftrightarrow$ truth; IoU / AP

HELD CASE · CONSTRUCTED PROBABILITY FIELD p_i is a teaching construction—not model output; blue cells satisfy $p_i \geq .50$.

truth G

0	1	1	0
0	1	1	0
0	1	1	0
0	0	0	0

**foreground p**

.10	.90	.90	.10
.60	.90	.90	.10
.10	.90	.30	.60
.10	.60	.05	.05

$$\sum_i p_i = 7.2, \quad \sum_i g_i = 6, \quad \sum_i p_i g_i = 4.8.$$

Binary cross-entropy scores each pixel locally

D DERIVATION

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

$$\ell_i = -[g_i \log p_i + (1 - g_i) \log(1 - p_i)].$$

FOREGROUND $\cdot g_i = 1$

penalty = $-\log p_i$
confident .90 $\rightarrow$.105
missed .30 $\rightarrow$ 1.204

BACKGROUND $\cdot g_i = 0$

penalty = $-\log(1 - p_i)$
safe .10 $\rightarrow$.105
false alarm .60 $\rightarrow$.916

CE asks whether each probability is calibrated to that pixel's label.

Group the held pixels to reproduce dense BCE

D DERIVATION

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

HELD CASE · SAME G, p Five foreground pixels have .90; one has .30. Background has $.60 \times 3, .10 \times 5, .05 \times 2$.

$$16\mathcal{L}_{\text{BCE}} = -10 \log .90 - \log .30 - 3 \log .40 - 2 \log .95$$

$$\approx 1.0536 + 1.2040 + 2.7489 + .1026 \approx 5.1090.$$

$$\mathcal{L}_{\text{BCE}} \approx \frac{5.109}{16} \approx .319.$$

The tensor keeps spatial axes separate from the class axis

D DERIVATION

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

Axis	Size	Meaning
h	H	vertical location
w	W	horizontal location
k	K	class score or probability

$$Z[h, w, :] \in \mathbb{R}^K \rightarrow p[h, w, :] \in [0, 1]^K.$$

Normalize over k ; reduce the loss over valid (h, w) locations.

TRAIN checkpoint: the labels supervise probabilities, not hard masks

QUESTION

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

Which object enters the derivative during training?

A

the thresholded mask P

B

the differentiable
probability p

C

only pixel accuracy

B. CE and Soft Dice operate on p ; thresholding belongs to INFER.

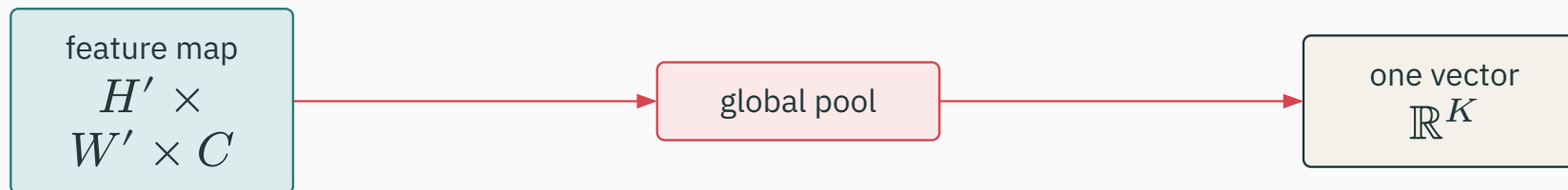
**TRAIN: preserve detail while
acquiring context**

A classification head destroys the layout segmentation needs

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP



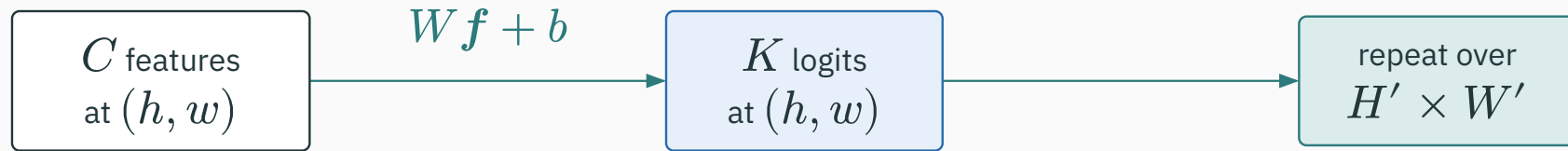
After locations collapse, different pixels cannot receive different labels.

A 1×1 convolution is the shared pixel classifier

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP



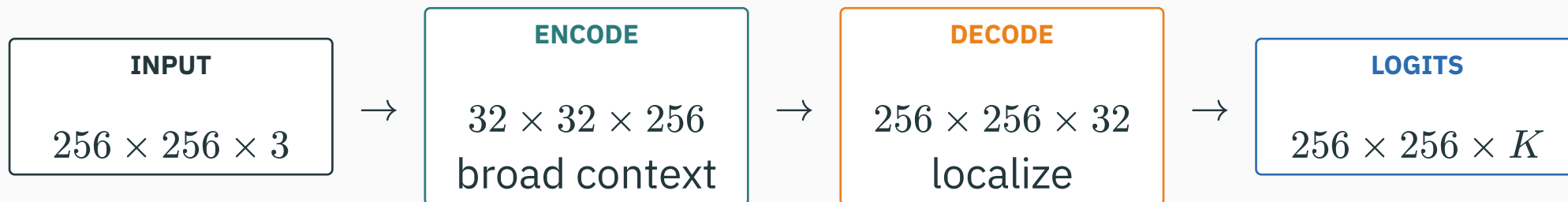
A 1×1 kernel mixes channels at one location. Earlier spatial convolutions already gathered neighborhood evidence.

Encoder–decoder networks exchange coordinates for context

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP



Downsample to understand what surrounds a pixel; upsample to decide which pixel.

A real boundary audit exposes what region overlap can hide

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

COMPUTED AUDIT

Region overlap can be strong while boundaries still shift

REGION · valid pixels only

BCE	0.1487
Soft Dice	0.5739
IoU	0.8108
Hard Dice	0.8955
Accuracy	0.9803

BOUNDARY · contour distances

Precision @ 5px	0.8098
Recall @ 5px	0.3027
Boundary F1 @ 5px	0.4407
ASSD (px)	6.940
HD95 (px)	25.495

Two contracts, one lesson

IoU/Dice reduce valid-region confusion. Boundary F1/ASSD/HD95 compare contours; the ignored GT band is the reference contour.

Boundary metrics are a diagnostic extension; they are not part of the deck's exact 4 x 4 numerical spine.

Does a larger grid recreate discarded detail?

QUESTION

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

A 2×2 feature map is repeated into a 4×4 grid. Did upsampling recover which source pixels created each value?

1	2
3	4



1	1	2	2
1	1	2	2
3	3	4	4
3	3	4	4

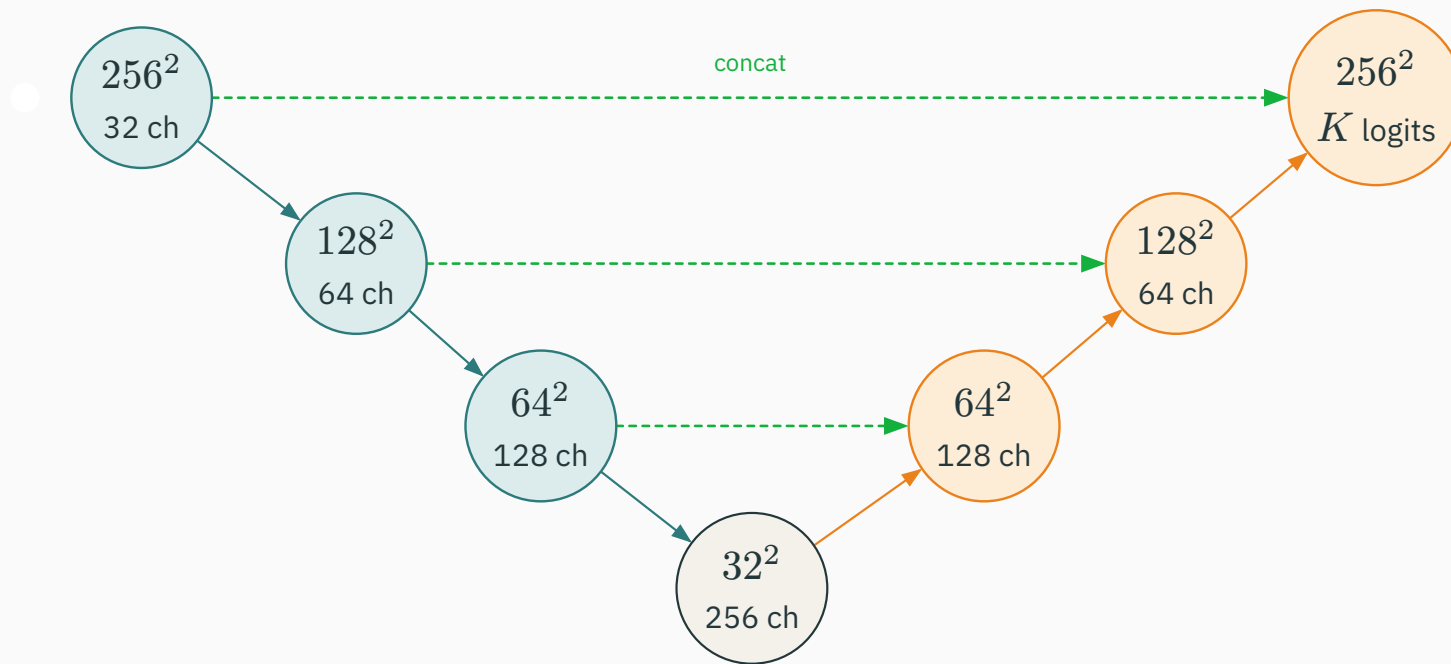
Commit before the next slide: resolution and information are not synonyms.

U-Net routes fine detail around the bottleneck

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP



Upsampling adds locations; skips restore evidence that the bottleneck discarded.

This ledger is a same-padding teaching U-Net—not the 2015 pixel sizes D DERIVATION

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

TOY LEDGER USED HERE

same-padding convolutions preserve each stage size; matching encoder and decoder tensors concatenate directly

ORIGINAL U-NET

valid convolutions shrink tensors; the encoder feature is copied **and cropped** before concatenation

$$U, E \in \mathbb{R}^{64 \times 64 \times 128} \rightarrow \text{concat}(U, E) \in \mathbb{R}^{64 \times 64 \times 256}.$$

The skip operation is concatenation in original U-Net; “skip connection” is not universally synonymous with addition.

Boundary failure has a smallest useful diagnostic

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

Symptom	Suspect	Smallest test
correct class; blurred edge	coarse stride or missing skip	overlay boundary; ablate one skip
checkerboard texture	transposed-conv overlap	compare resize+conv at same output size
shifted instance mask	RoI quantization	compare RoIPool with RoIAlign

Inspect spatial alignment before changing the loss.

Architecture checkpoint: context and coordinates play different roles QUESTION

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

If the deep path is removed but every high-resolution skip remains, what is most at risk?

A

semantic meaning

B

number of pixels

C

softmax normalization

A. A crisp local edge does not by itself say whether the region is road, cat, or shadow.

**INFER makes a mask; EVAL
measures its agreement**

TRAIN ends with probabilities; INFER begins without labels

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

$$G, p \rightarrow \mathcal{L}_{\text{BCE}}, \mathcal{L}_{\text{Dice}}$$

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

$$p, \tau \rightarrow P; \text{no truth and no loss}$$

The same p crosses a clock boundary; its meaning does not change.

Commit: threshold the held probabilities at .50

QUESTION

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

HELD CASE · DECISION RULE $P_i = 1[p_i \geq .50]$. Equality is foreground by convention.

truth G

0	1	1	0
0	1	1	0
0	1	1	0
0	0	0	0

$\rightarrow$

foreground p

.10	.90	.90	.10
.60	.90	.90	.10
.10	.90	.30	.60
.10	.60	.05	.05

Before revealing P , predict the four counts: TP, FP, FN, TN.

Thresholding yields this exact hard mask

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

truth G

0	1	1	0
0	1	1	0
0	1	1	0
0	0	0	0



foreground p

.10	.90	.90	.10
.60	.90	.90	.10
.10	.90	.30	.60
.10	.60	.05	.05

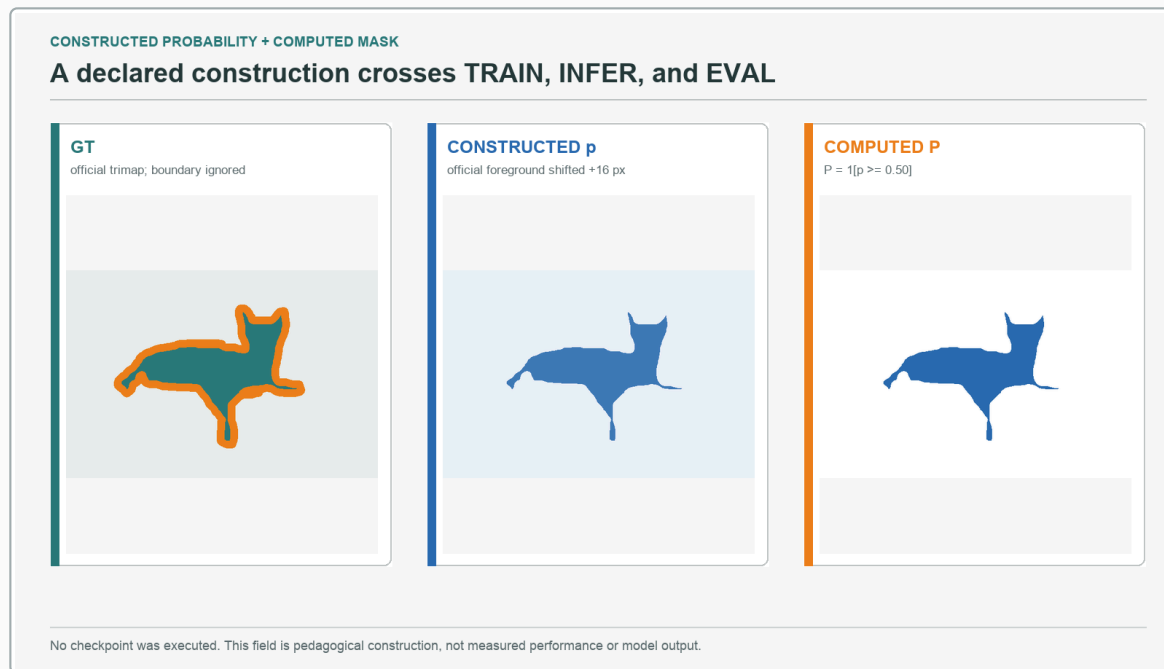


hard P

0	1	1	0
1	1	1	0
0	1	0	1
0	1	0	0

Thresholding is nondifferentiable; it belongs to inference or metric computation, not the training path.

The same clocks operate on actual pixels



GT is official; p is the declared +16 px construction; $P = 1[p \geq .50]$ is computed.
No checkpoint was executed.

Count errors before naming a metric

TRAIN probability $\leftrightarrow$ truth; CE / Dice**INFER** probability $\rightarrow$ label / object masks**EVAL** hard output $\leftrightarrow$ truth; IoU / APtruth G

0	1	1	0
0	1	1	0
0	1	1	0
0	0	0	0

prediction P

0	1	1	0
1	1	1	0
0	1	0	1
0	1	0	0

$$TP = 5 \quad FP = 3 \quad FN = 1 \quad TN = 7$$

Every hard metric below is a different reduction of the same four counts.

Pixel accuracy can reward an empty foreground prediction

D DERIVATION

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

HELD CASE · HELD MASK accuracy = $\frac{TP+TN}{16} = \frac{5+7}{16} = .750$.

Now imagine a 100×100 image with only 100 tumor pixels.

predict background everywhere $\rightarrow$ accuracy = $\frac{9900}{10000} = 99\%$.

tumor recall = $\frac{0}{100} = 0$.

Accuracy weights the easy background by its frequency; it may not reflect the foreground task.

IoU removes true background from the denominator

D DERIVATION

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

$$\text{IoU} = \frac{|P \cap G|}{|P \cup G|} = \frac{\text{TP}}{\text{TP} + \text{FP} + \text{FN}}.$$

HELD CASE · SAME COUNTS $|P \cap G| = 5$ and $|P \cup G| = 5 + 3 + 1 = 9$.

$$\text{IoU} = \frac{5}{9} \approx .556.$$

Seven true-background pixels cannot inflate foreground IoU.

Real pixels reveal where every region count lives

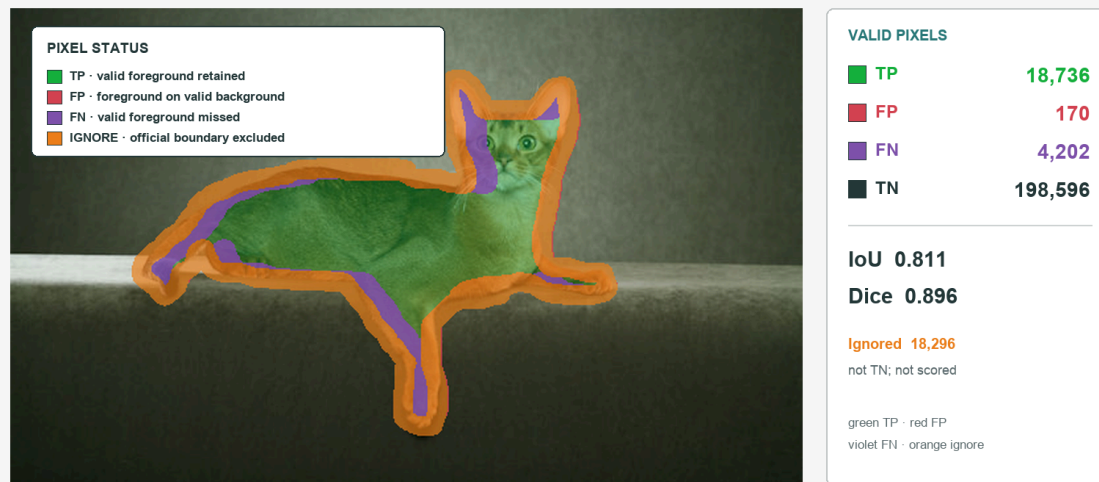
TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

COMPUTED OVERLAY

Region counts become visible on the real photograph



GT and photo are real. Probability, prediction, error colors, counts, and metrics are computed from the declared construction.

Dice gives the intersection twice the weight

D DERIVATION

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

$$\text{Dice} = \frac{2|P \cap G|}{|P| + |G|} = \frac{2\text{TP}}{2\text{TP} + \text{FP} + \text{FN}}.$$

HELD CASE · SAME COUNTS $|P| = 8$, $|G| = 6$, and $|P \cap G| = 5$.

$$\text{Dice} = \frac{10}{14} \approx .714.$$

For one binary mask, Dice is the foreground F1 score.

Dice is a monotone transform of IoU for one mask pair

D DERIVATION

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

$$\text{Dice} = \frac{2\text{IoU}}{1+\text{IoU}}, \quad \text{IoU} = \frac{\text{Dice}}{2-\text{Dice}}.$$

$$\frac{2(\frac{5}{9})}{1 + \frac{5}{9}} = \frac{10}{14} \approx .714.$$

Qualifier: this preserves ranking for a single class and mask pair. Macro-averaging nonlinear scores over classes or images can change a model ranking.

Soft Dice scores the held probabilities before thresholding

D DERIVATION

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

$$\text{SoftDice} = \frac{2 \sum_i p_i g_i + \varepsilon}{\sum_i p_i + \sum_i g_i + \varepsilon}.$$

HELD CASE · SAME G, p $\sum p = 7.2, \sum g = 6, \sum pg = 4.8$; omit tiny ε in the displayed arithmetic.

$$\text{SoftDice} = \frac{9.6}{13.2} \approx .727.$$

$$\mathcal{L}_{\text{Dice}} = 1 - .727 = .273.$$

CE and Soft Dice expose different failure surfaces

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

LOCAL CALIBRATION · CE

scores every valid pixel; strong gradients identify which probabilities are wrong

REGION OVERLAP · DICE

couples foreground pixels through one region denominator; resists background domination

$$\mathcal{L} = \mathcal{L}_{\text{CE}} + \lambda \mathcal{L}_{\text{Dice}}.$$

Combining them does not make either definition optional: declare λ , reduction, and empty-mask policy.

Semantic mIoU averages declared per-class IoUs

D DERIVATION

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

For each class k ,

$$\text{IoU}_k = \frac{\text{TP}_k}{\text{TP}_k + \text{FP}_k + \text{FN}_k}, \quad \text{mIoU} = \frac{1}{|K_{\text{report}}|} \sum_{k \in K_{\text{report}}} \text{IoU}_k.$$

CLASSES

VOID

REDUCTION

include or exclude
background?

remove ignored pixels
from all counts

image mean or dataset-
level confusion?

“mIoU” is incomplete until its class set and aggregation rule are named.

Symptom	Suspect	Smallest test
99% accuracy; IoU = 0	rare foreground	foreground recall; class counts
loss falls; mIoU flat	threshold or class collapse	histogram p by class; confusion
IoU .811; boundary F1 .441	spatial displacement	overlay FP/FN; inspect contour distances

Region and boundary metrics interrogate different geometry; choose the diagnostic on EVAL, then inspect actual pixels.

**One semantic region is not yet
two objects**

A generated touching-object scene isolates the identity question

generated concept scene



constructed union

0	1	1	0
0	1	1	0
0	1	1	0
0	0	0	0

one connected component

constructed identities

0	1	2	0
0	1	2	0
0	1	2	0
0	0	0	0

id 1 vs id 2

Oxford supplies real semantic GT; this explicitly generated scene demonstrates why a semantic union cannot count touching instances.

Instance segmentation returns a scored set—not one class grid

D DERIVATION

TRAIN probability $\leftrightarrow$ truth; CE / Dice

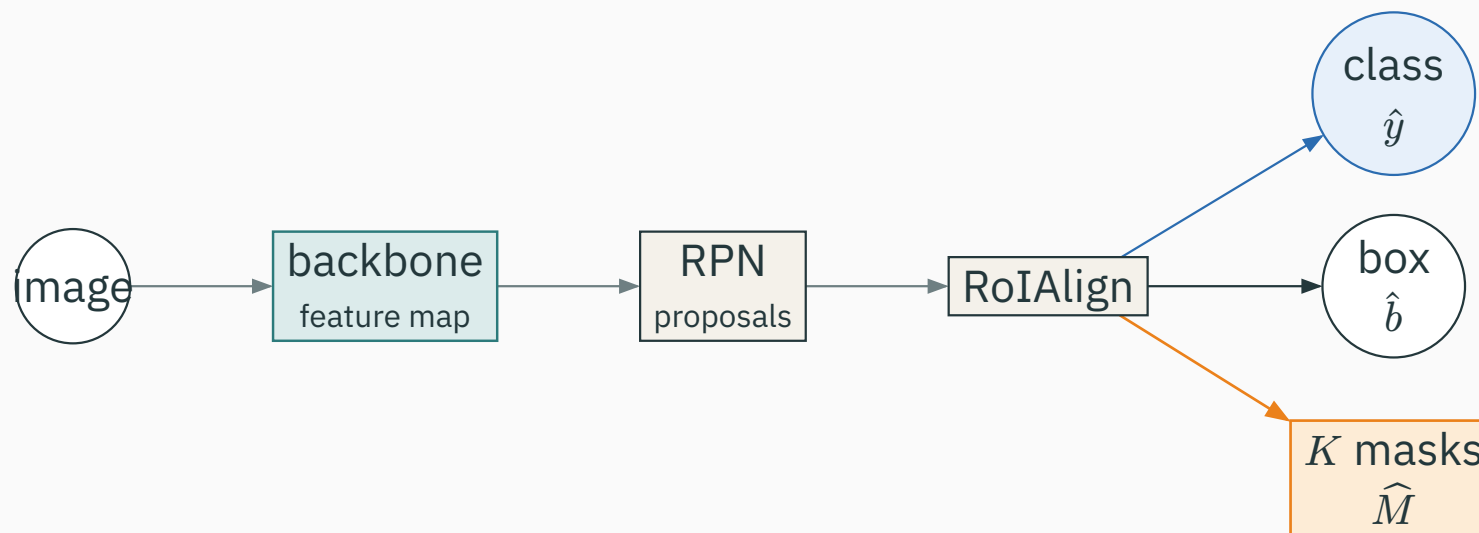
INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

$$x \rightarrow \left\{ \left(\hat{y}_j, \hat{b}_j, \hat{m}_j, \hat{s}_j \right) \right\}_{j=1}^N.$$

Symbol	Shape	Meaning
$\hat{y}_j$	scalar	object class
$\hat{b}_j$	four coordinates	box that places the mask
$\hat{m}_j$	$m \times m$ then resized	shape in the RoI frame
$\hat{s}_j$	scalar	ranking confidence

Mask R-CNN adds a mask branch parallel to class and box



RPN = Region Proposal Network; RoI = Region of Interest.

Detection proposes **which object**; the mask branch predicts **which RoI pixels**.

Separate the RPN loss from the second-stage RoI loss

D DERIVATION

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

STAGE 1 · RPN

$\mathcal{L}_{\text{RPN-obj}}$ on sampled anchors
 $\mathcal{L}_{\text{RPN-box}}$ on positive anchors

STAGE 2 · SAMPLED ROIS

$\mathcal{L}_{\text{cls}}$ on positive + negative RoIs
 $\mathcal{L}_{\text{box}} + \mathcal{L}_{\text{mask}}$ on positive RoIs

$$\mathcal{L}_{\text{total}} = \mathcal{L}_{\text{RPN}} + \mathcal{L}_{\text{RoI}}.$$

The familiar three-term Mask R-CNN equation is the **second-stage RoI loss**, not the whole RPN-plus-RoI total.

Original Mask R-CNN predicts K independent binary masks

D DERIVATION

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

$$\hat{M} \in \mathbb{R}^{K \times m \times m}, \quad \hat{M}_{k,u,v} = \sigma(a_{k,u,v}).$$

For a training RoI with ground-truth class k , only channel k contributes to $\mathcal{L}_{\text{mask}}$.

CLASS HEAD

softmax chooses the object category

MASK CHANNEL k

sigmoid makes foreground/
background decisions inside that RoI

A class-agnostic one-mask head is a valid later variant; it is not the original K -mask formulation.

Calculate one selected-channel mask BCE

D DERIVATION

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

$$M = \begin{pmatrix} 1 & 1 \\ 0 & 0 \end{pmatrix}, \quad \hat{M} = \begin{pmatrix} .8 & .6 \\ .3 & .1 \end{pmatrix}.$$

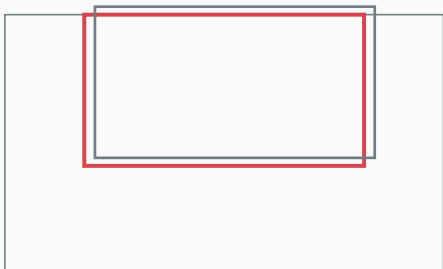
$$\mathcal{L}_{\text{mask}} = -\frac{1}{4}[\log .8 + \log .6 + \log(1 - .3) + \log(1 - .1)]$$

$$= \frac{1}{4} [.223 + .511 + .357 + .105] \approx .299.$$

Pixels in this RoI and the selected ground-truth class channel define the mask loss.

ROIPOOL

round RoI and bin coordinates;
convenient, but spatially quantized



ROIALIGN

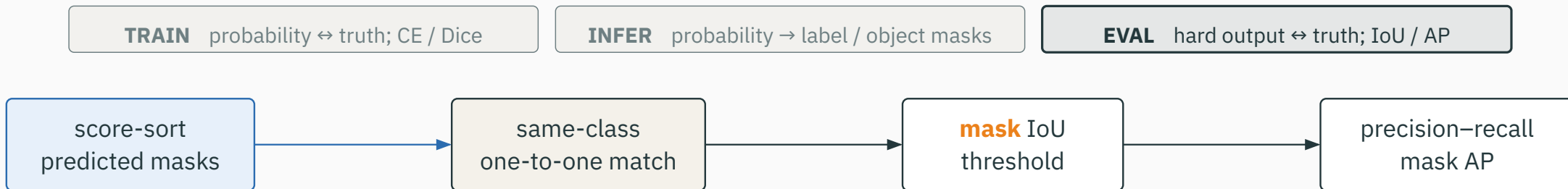
keep fractional coordinates; sample
neighboring features bilinearly



The Mask R-CNN paper reports box-AP gains too; the alignment benefit is larger for pixel-accurate masks.

Instance evaluation reuses L10 AP with mask IoU

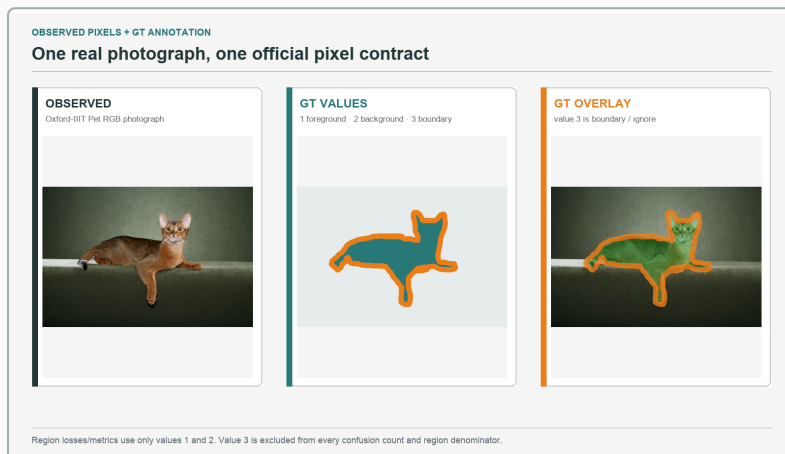
D DERIVATION



Box AP asks whether rectangles overlap; mask AP asks whether object pixels overlap.

Policies, alignment, and panoptic output

Void pixels leave both the loss sum and denominator

TRAIN probability $\leftrightarrow$ truth; CE / Dice**INFER** probability $\rightarrow$ label / object masks**EVAL** hard output $\leftrightarrow$ truth; IoU / AP

Let V be the valid labeled pixels:

$$\mathcal{L}_{\text{CE}} = -\frac{1}{|V|} \sum_{i \in V} \log p_{i, y_i}.$$

Oxford values 1/2 are valid: $|V| = 221,704$. Value 3 contributes 18,296 ignored boundary pixels.

Remove ignored pixels from CE, confusion counts, IoU, and Dice—not only from the numerator.

Empty masks and class averaging need a written policy

D DERIVATION

TRAIN probability $\leftrightarrow$ truth; CE / Dice

INFER probability $\rightarrow$ label / object masks

EVAL hard output $\leftrightarrow$ truth; IoU / AP

EMPTY-EMPTY

raw Dice is $\frac{0}{0}$; choose smoothing, exclusion, or a declared perfect score

MULTI-CLASS

compute per class; declare background inclusion, empty-class handling, and macro or frequency weighting

Also declare whether scores are averaged per image or accumulated from a dataset-level confusion matrix.

The reduction rule is part of the metric—not an implementation footnote.

Upsampling restores resolution by different mechanisms

Method	What creates locations?	What learns?
nearest / bilinear	fixed interpolation	following convolution refines
transposed convolution	learned overlap pattern	kernel and mixing jointly
skip + resize	decoder plus aligned encoder detail	fusion after concatenation

Every route needs evidence to refine the larger grid; interpolation alone cannot reverse pooling.

Bilinear interpolation is a weighted four-neighbor sum

Suppose a RoIAlign sample lies halfway in both directions between

$$\begin{pmatrix} 1 & 3 \\ 5 & 7 \end{pmatrix}.$$

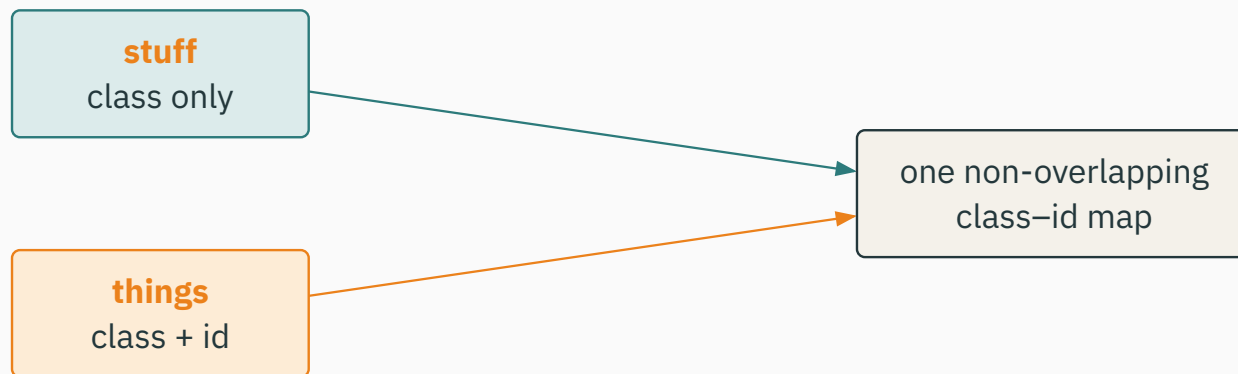
All four weights equal .25:

$$f = .25(1) + .25(3) + .25(5) + .25(7) = 4.$$

Moving the fractional sample changes the weights smoothly; no coordinate is rounded.

Precise sampling preserves alignment; it does not invent image detail.

Panoptic output assigns one class–identity record per pixel



Overlapping instance hypotheses must be fused; datasets may reserve a void label for uncertain pixels.

Choose the output and metric from the question

Need	Output	Evaluation
road area	$H \times W$ semantic classes	per-class IoU / mIoU
count touching trees	scored set of instance masks	mask AP under mask IoU
whole labeled scene	$H \times W$ class-id map	PQ or declared panoptic protocol

Architecture names follow the output contract; they do not define the task.

Primary papers establish the architectural claims

FCN + U-NET

Long et al. · pixels-to-pixels; Ronneberger et al. · context plus localization

PANOPTIC + COMPANION

Kirillov et al. · stuff + things; exact L10 Colab · toy and real ledgers

REPRODUCIBILITY

shared/vision-evidence/oxford-iiit-pet/l10/ · builder + manifest + arrays + metric ledger

OBSERVED = photo · GT = trimap · CONSTRUCTED = orchard / toy p / shifted p · COMPUTED = masks, errors, metrics

V-NET + MASK R-CNN

Milletari et al. · Dice objective; He et al. · mask branch + RoIAlign

REAL EVIDENCE

Oxford-IIIT Pet · Abyssinian_1 · official trimap · CC BY-SA 4.0

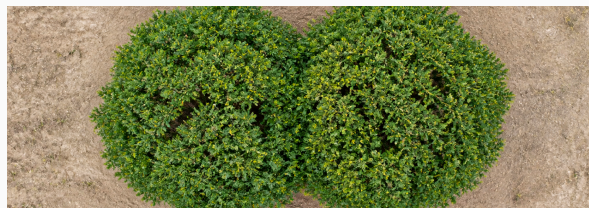
HONEST CLAIM

model output = none · measured performance = none · +16 px probability field = constructed

Variants change mechanisms, not the contract

Variant	What changes	What stays explicit
class-agnostic mask head	$1 \times m \times m$ rather than $K \times m \times m$	object class comes from class head
dilated / pyramid decoder	context without identical downsampling	output grid and pixel loss
query-based mask sets	learned queries replace RPN RoIs	one-to-one identity and mask metric
panoptic fusion	resolve overlaps and stuff regions	one class-identity assignment per pixel

Do not turn variants into a catalogue: ask which output, alignment, loss, or matching assumption changed.



HELD CASE · SAME GENERATED CONCEPT Report total canopy area **and** count two touching trees.

A

global class

B

class + box

C

semantic map

D

separate object masks

D. Instance masks preserve the two identities; their union gives total area. If every background/stuff pixel also needs a class, use a panoptic class–identity map.

SEGMENTATION CHECKLIST

1. What is the output structure?
2. Which axis holds classes?
3. What enters TRAIN loss?
4. Which INFER threshold/fusion applies?
5. Which EVAL classes, masks, and reductions define the score?

PUBLIC L12 · NEXT TOKEN

Pixel prediction shared one classifier over spatial positions (h, w) .

Next-token prediction shares one classifier over sequence positions t .

Pixel classes become vocabulary tokens; **causality** becomes the new constraint.

A dense model is trustworthy only when output support, three clocks, alignment, and metric protocol agree.