

ECE 60146 — Homework 9

Transformers on Translation, CIFAR-10, and Food-101

Austin Toro

Spring 2026

Contents

1	Task 1: Seq2seq Translation Warm-up (15 points)	3
1.1	(a) Code: Script used to run TransformerFG	3
1.2	(b) Figure: Training loss curve over 4 epochs	6
1.3	(c) Text: 20 translations with correctness labels and accuracy fraction	6
1.4	(d) Text: Warm-up phase observations	7
1.5	(e) Text: Self-attention computation explanation	8
1.6	(f) Text: Multi-headed attention explanation	8
1.7	(g) Text: TransformerFG vs. TransformerPreLN comparison	8
2	Task 2: Vision Transformer Warm-up on CIFAR-10 (15 points)	9
2.1	(a) Code: Script used to run visTransformer	9
2.2	(b) Figure: Training loss curve over 40 epochs	12
2.3	(c) Table: 10×10 confusion matrix on test set	13
2.4	(d) Text: Overall test accuracy	14
2.5	(e) Text: Patch embedding explanation	14
2.6	(f) Text: Class token explanation	14
2.7	(g) Text: Positional encoding explanation	14
3	Task 3: Food-101 From Scratch + Attention Visualization (25 points)	15
3.1	(a) Code: Adapted visTransformer for Food-101 subset	15
3.2	(b) Figure: Training loss curve	20
3.3	(c) Text: Test accuracy and training observations	20
3.4	(d) Code: Modified AttentionHead, SelfAttention, BasicEncoder, and MasterEncoder	20
3.5	(e) Text: Description of attention extraction modifications	23
3.6	(f) Figure: Part A — 4 images with full heatmap and spatial overlay	25
3.7	(g) Text: Part A — Per-image attention analysis	26
3.8	(h) Figure: Part B — 2 pairs with side-by-side overlays	26
3.9	(i) Text: Part B — Comparative analysis per pair	26

4	Task 4: Pretrained ViT Fine-tuning on Food-101 Subset (45 points)	28
4.1	(a) Code: Fine-tuning script for ViT-Base on Food-101 subset	28
4.2	(b) Figure: Training loss and validation accuracy curves	32
4.3	(c) Text: Test accuracy and parameter count	32
4.4	(d) Figure: 4×4 image grid with predictions	33
4.5	(e) Table: Two-model comparison table	34
4.6	(f) Text: Accuracy gap discussion	34
4.7	(g) Text: Model capacity and data requirements discussion	34
4.8	(h) Text: Connection to language model pretraining	34

1 Task 1: Seq2seq Translation Warm-up (15 points)

1.1 (a) Code: Script used to run TransformerFG

```
1 class InstrumentedTransformerFG(TransformerFG):
2     """TransformerFG wrapper that records metrics needed for the
3         report."""
4
5     def train_and_record(
6         self,
7         master_encoder: torch.nn.Module,
8         master_decoder: torch.nn.Module,
9         checkpoint_dir: str | Path,
10    ) -> dict[str, Any]:
11        checkpoint_dir = ensure_dir(checkpoint_dir)
12        master_encoder.to(self.dl_studio.device)
13        master_decoder.to(self.dl_studio.device)
14        embeddings_generator_en = self.EmbeddingsGenerator(self, "en"
15            , self.embedding_size).to(self.dl_studio.device)
16        embeddings_generator_es = self.EmbeddingsGenerator(self, "es"
17            , self.embedding_size).to(self.dl_studio.device)
18        optimizer_params = self.optimizer_params
19        beta1 = optimizer_params["beta1"]
20        beta2 = optimizer_params["beta2"]
21        epsilon = optimizer_params["epsilon"]
22        master_encoder_optimizer = self.ScheduledOptim(
23            optim.Adam(master_encoder.parameters(), betas=(beta1,
24                beta2), eps=epsilon),
25            lr_mul=2,
26            d_model=self.embedding_size,
27            n_warmup_steps=self.num_warmup_steps,
28        )
29        master_decoder_optimizer = self.ScheduledOptim(
30            optim.Adam(master_decoder.parameters(), betas=(beta1,
31                beta2), eps=epsilon),
32            lr_mul=2,
33            d_model=self.embedding_size,
34            n_warmup_steps=self.num_warmup_steps,
35        )
36        criterion = nn.NLLLoss()
37        num_sentence_pairs = len(self.training_corpus)
38        batch_size = self.dl_studio.batch_size
39        max_iters = math.ceil(num_sentence_pairs / batch_size)
40        history: dict[str, Any] = {
41            "loss_points": [],
42            "learning_rate_points": [],
43            "iterations_logged": [],
```

```

39         "epoch_average_losses": [],
40     }
41     start_time = time.perf_counter()
42     for epoch in range(self.dl_studio.epochs):
43         random.shuffle(self.training_corpus)
44         running_loss = 0.0
45         epoch_running_loss = 0.0
46         epoch_examples = 0
47         for iteration in range(max_iters):
48             left = iteration * batch_size
49             right = min((iteration + 1) * batch_size,
50                          num_sentence_pairs)
51             batched_pairs = self.training_corpus[left:right]
52             source_sentences = [pair[0] for pair in batched_pairs]
53             target_sentences = [pair[1] for pair in batched_pairs]
54             en_tensor = self.sentence_with_words_to_ints(
55                 source_sentences, "en").to(self.dl_studio.device)
56             es_tensor = self.sentence_with_words_to_ints(
57                 target_sentences, "es").to(self.dl_studio.device)
58             en_sentence_tensor = embeddings_generator_en(
59                 en_tensor).float()
60             es_sentence_tensor = embeddings_generator_es(
61                 es_tensor).float()
62             master_encoder_optimizer.zero_grad()
63             master_decoder_optimizer.zero_grad()
64             master_encoder_output = master_encoder(
65                 en_sentence_tensor)
66             predicted_word_logprobs, _ = master_decoder(
67                 es_sentence_tensor, master_encoder_output)
68             loss = torch.tensor(0.0, device=self.dl_studio.device)
69             for column in range(es_tensor.shape[1]):
70                 loss = loss + criterion(predicted_word_logprobs
71                                        [:, column], es_tensor[:, column])
72             loss.backward()
73             master_encoder_optimizer.step_and_update_lr()
74             master_decoder_optimizer.step_and_update_lr()
75
76             normalized_loss = loss.item() / es_tensor.shape[0]
77             running_loss += normalized_loss
78             epoch_running_loss += normalized_loss * es_tensor.
79                 shape[0]
80             epoch_examples += es_tensor.shape[0]
81
82             if iteration % 200 == 199:

```

```

74         avg_loss = running_loss / 200.0
75         running_loss = 0.0
76         current_lr = master_encoder_optimizer._optimizer.
            param_groups[0]["lr"]
77         elapsed_seconds = time.perf_counter() -
            start_time
78         print(
79             "[epoch:%2d/%d  iter:%4d  elapsed_time:%5d
            secs] loss: %.4f  lr: %.8f"
80             % (
81                 epoch + 1,
82                 self.dl_studio.epochs,
83                 iteration + 1,
84                 elapsed_seconds,
85                 avg_loss,
86                 current_lr,
87             )
88         )
89         history["loss_points"].append(avg_loss)
90         history["learning_rate_points"].append(current_lr)
91         history["iterations_logged"].append((epoch *
            max_iters) + iteration + 1)
92
93         epoch_average_loss = epoch_running_loss / max(
            epoch_examples, 1)
94         history["epoch_average_losses"].append(epoch_average_loss)
95         print(f"Finished epoch {epoch + 1}/{self.dl_studio.epochs}
            with mean loss {epoch_average_loss:.4f}")
96
97         self.save_encoder(master_encoder)
98         self.save_decoder(master_decoder)
99         self.save_embeddings_generator_en(embeddings_generator_en)
100        self.save_embeddings_generator_es(embeddings_generator_es)
101        history["elapsed_seconds"] = time.perf_counter() - start_time
102        return history

```

Listing 1: Instrumented TransformerFG training wrapper adapted from DLStudio’s seq2seq example and training loop

1.2 (b) Figure: Training loss curve over 4 epochs

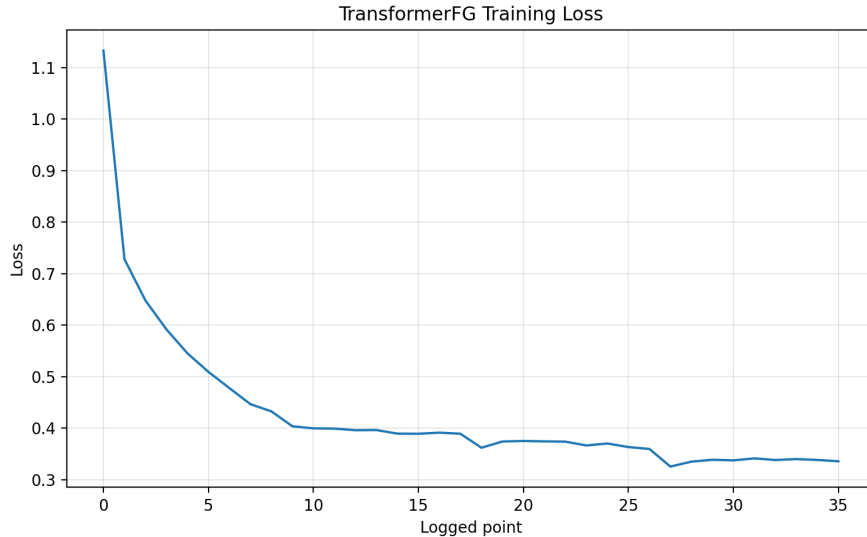


Figure 1: Task 1 TransformerFG training loss over the 4-epoch run. Logged once every 200 iterations.

1.3 (c) Text: 20 translations with correctness labels and accuracy fraction

Exact-match accuracy was 1/20. However, 15 additional outputs were at least partially semantically aligned with the target sentence, usually by preserving the main noun, verb, or sentence frame while dropping details or repeating words.

Table 1: Task 1 sampled translations with coarse correctness labels.

#	English source	Spanish target	Model output	Label
1	SOS my down jacket keeps me quite warm EOS	SOS mi chaqueta de plumas me mantiene bastante calentito EOS	EOS mi me chaqueta chaqueta mantiene bastante bastante EOS EOS	semantically correct
2	SOS stop resisting EOS	SOS deja de resistirte EOS	EOS deja de EOS EOS EOS EOS EOS EOS EOS	semantically correct
3	SOS it is a beautiful day is not it EOS	SOS bonito día verdad EOS	EOS es un hermoso EOS EOS EOS EOS EOS EOS	wrong
4	SOS i never imagined meeting you here EOS	SOS nunca imaginé encontrarte aquí EOS	EOS nunca he reunión aquí aquí aquí EOS EOS EOS	semantically correct
5	SOS tom took mary home EOS	SOS tom llevó a mary a casa EOS	EOS tom se a a a a EOS EOS EOS	wrong
6	SOS i tried to get back to sleep EOS	SOS intenté volverme a dormir EOS	EOS intenté de a EOS EOS EOS EOS EOS EOS	semantically correct

#	English source	Spanish target	Model output	Label
7	SOS something remarkable happened EOS	SOS ocurrió algo increíble EOS	EOS algo algo EOS EOS EOS EOS EOS EOS EOS	wrong
8	SOS i am happy and satisfied EOS	SOS estoy feliz y satisfecho EOS	EOS estoy estoy satisfecho satisfecho EOS EOS EOS EOS EOS EOS	semantically correct
9	SOS my apartment is near EOS	SOS mi apartamento está cerca EOS	EOS mi apartamento está cerca cerca EOS EOS EOS EOS	semantically correct
10	SOS we eat soup with a spoon EOS	SOS comemos la sopa con cuchara EOS	EOS comemos sopa con una EOS EOS EOS EOS EOS EOS	semantically correct
11	SOS i do not have to answer that question EOS	SOS no tengo que contestar a esa pregunta EOS	EOS no tengo que que esa esa EOS EOS EOS EOS	semantically correct
12	SOS this book is new EOS	SOS este libro es nuevo EOS	EOS este libro es nuevo EOS EOS EOS EOS EOS	correct
13	SOS you can talk to them EOS	SOS puedes hablar con ellos EOS	EOS puedes hablar con EOS EOS EOS EOS EOS EOS	semantically correct
14	SOS she is able to sing very well EOS	SOS ella puede cantar muy bien EOS	EOS ella es cantar muy muy EOS EOS EOS EOS	semantically correct
15	SOS tom also has plans to go there EOS	SOS tom también planea ir allí EOS	EOS tom también tiene ir ir EOS EOS EOS EOS	semantically correct
16	SOS even tom is surprised that mary lied EOS	SOS incluso a tom le sorprendió que maría mintiera EOS	EOS ni es sorprendió sorprendió sorprendió a mary EOS EOS EOS	wrong
17	SOS i am in dire need of money EOS	SOS estoy mal de dinero EOS	EOS estoy en de de dinero EOS EOS EOS EOS EOS	semantically correct
18	SOS this is the real world EOS	SOS este es el mundo real EOS	EOS éste es el el mundo EOS EOS EOS EOS EOS	semantically correct
19	SOS i work in boston EOS	SOS trabajo en boston EOS	EOS trabajo en en EOS EOS EOS EOS EOS EOS	semantically correct
20	SOS he began to run EOS	SOS él empezó a correr EOS	EOS empezó empezó a EOS EOS EOS EOS EOS EOS	semantically correct

1.4 (d) Text: Warm-up phase observations

The loss shows the expected warm-up dynamics. During the first 4,000 optimization steps the scheduled learning rate ramps from roughly 9.9×10^{-5} to 1.98×10^{-3} , and the loss drops rapidly from about 1.13 at iteration 200 to about 0.37 near the end of warm-up. After step 4,000 the schedule begins decaying the learning rate and the loss continues to improve, but much more gradually, ending near 0.3366 by epoch 4.

Qualitatively, the model clearly learned common sentence structure and many high-frequency Spanish content words within only four epochs, but it still repeated tokens,

dropped function words, and occasionally collapsed to short partial phrases. This is the expected behavior as there are not enough runs to produce a strong translator.

1.5 (e) Text: Self-attention computation explanation

If the input sentence contains N_w tokens and each token embedding has width M , then the learned query and key tensors are $Q, K \in \mathbb{R}^{N_w \times M}$. The matrix product $QK^\top$ therefore has shape $N_w \times N_w$: each entry compares one token’s query against another token’s key. After row-wise **Softmax**, row i becomes a probability distribution over all source positions, so it tells us how strongly token i attends to every token in the sentence, including itself.

The division by $\sqrt{M}$ prevents the raw dot products from growing too large when the embedding dimension is large. Without that scaling, the logits entering the softmax would have high variance, the softmax would become overly peaked early in training, and the gradients would become unstable or very small. The scaling keeps attention scores in a numerically reasonable range.

1.6 (f) Text: Multi-headed attention explanation

With embedding size 256 and 4 attention heads, each head receives a slice of size

$$\text{sqkv} = 256/4 = 64.$$

Each head therefore learns its own query, key, and value projections in a 64-dimensional subspace, and the outputs are concatenated afterward.

Multiple heads are better than a single full-width head because different heads can specialize in different relationships at the same time. One head may learn local syntactic alignment, another may focus on subject–verb agreement, and another may lock onto longer-range semantic dependencies. A single head with all 256 channels would have to compress all of those roles into one attention pattern.

1.7 (g) Text: TransformerFG vs. TransformerPreLN comparison

In TransformerFG, LayerNorm is applied after each sublayer’s residual addition, so the pattern is “self-attention $\rightarrow$ residual add $\rightarrow$ LayerNorm” and then “FFN $\rightarrow$ residual add $\rightarrow$ LayerNorm.” In TransformerPreLN, LayerNorm is moved before the self-attention and FFN sublayers, so the residual path itself is more direct.

That placement difference explains the warm-up requirement. TransformerFG is a post-norm design, which is less stable during the earliest optimization steps because gradients must pass through more poorly conditioned transformations before the model has learned useful internal scales. TransformerPreLN keeps the residual pathway cleaner, so gradient flow is more stable from the start and the model can usually train with a fixed learning rate instead of a fragile warm-up schedule.

2 Task 2: Vision Transformer Warm-up on CIFAR-10 (15 points)

2.1 (a) Code: Script used to run visTransformer

```
1 def main() -> None:
2     """Train the CIFAR-10 visTransformer and evaluate the best
3         checkpoint."""
4
5     args = parse_args()
6     setup_seed(args.seed)
7     results_dir = ensure_dir(PROJECT_ROOT / args.results_dir)
8     figures_dir = ensure_dir(PROJECT_ROOT / args.figures_dir)
9     checkpoint_dir = ensure_dir(results_dir / "checkpoints_visTrans")
10
11     dls = DLStudio(
12         dataroot=str(PROJECT_ROOT / args.data_root),
13         image_size=(32, 32),
14         path_saved_model=str(results_dir / "saved_visTran_model"),
15         learning_rate=args.learning_rate,
16         epochs=args.epochs,
17         batch_size=args.batch_size,
18         classes=CIFAR10_CLASSES,
19         use_gpu=True,
20     )
21     device = dls.device
22
23     model = visTransformer(
24         dl_studio=dls,
25         patch_size=(8, 8),
26         embedding_size=args.embedding_size,
27         num_basic_encoders=args.num_basic_encoders,
28         num_attn_heads=args.num_attn_heads,
29         save_checkpoints=False,
30     )
31     model.to(device)
32     patch_embedding_generator = model.PatchEmbeddingGenerator(model).
33         to(device)
34     train_loader, test_loader = build_loaders(args.batch_size,
35         PROJECT_ROOT / args.data_root)
36
37     optimizer = torch.optim.Adam(model.parameters(), lr=args.
38         learning_rate, betas=(0.9, 0.99))
39     patch_optimizer = torch.optim.Adam(patch_embedding_generator.
40         parameters(), lr=args.learning_rate, betas=(0.9, 0.99))
41     criterion = nn.CrossEntropyLoss()
```

```

38 training_loss_history: list[float] = []
39 checkpoint_scores: list[dict[str, float]] = []
40 start_time = time.perf_counter()
41
42 for epoch in range(1, args.epochs + 1):
43     model.train()
44     patch_embedding_generator.train()
45     running_loss = 0.0
46     example_count = 0
47     for images, labels in train_loader:
48         images = images.to(device, non_blocking=True)
49         labels = labels.to(device, non_blocking=True)
50         optimizer.zero_grad()
51         patch_optimizer.zero_grad()
52         tokens = make_tokens(model, patch_embedding_generator,
53                               images)
54         logits = model(tokens)
55         loss = criterion(logits, labels)
56         loss.backward()
57         optimizer.step()
58         patch_optimizer.step()
59         running_loss += loss.item() * labels.shape[0]
60         example_count += labels.shape[0]
61
62     mean_train_loss = running_loss / max(example_count, 1)
63     training_loss_history.append(mean_train_loss)
64     elapsed = time.perf_counter() - start_time
65     print(f"epoch {epoch:02d}/{args.epochs} train_loss={
66           mean_train_loss:.4f} elapsed={elapsed:.1f}s")
67
68     if epoch % args.checkpoint_freq == 0:
69         save_checkpoint(model, patch_embedding_generator,
70                         checkpoint_dir, epoch)
71         checkpoint_eval = evaluate(model,
72                                    patch_embedding_generator, test_loader, device)
73         checkpoint_scores.append({"epoch": epoch, "accuracy":
74                                   checkpoint_eval["accuracy"]})
75         print(f"checkpoint epoch {epoch}: accuracy={
76               checkpoint_eval['accuracy']:.4%}")
77
78 if not checkpoint_scores:
79     checkpoint_scores.append({"epoch": args.epochs, "accuracy":
80                               evaluate(model, patch_embedding_generator, test_loader,
81                                         device)["accuracy"]})
82     save_checkpoint(model, patch_embedding_generator,
83                     checkpoint_dir, args.epochs)

```

```

76     best_checkpoint = max(checkpoint_scores, key=lambda item: item["
77         accuracy"])
78     load_checkpoint(model, patch_embedding_generator, checkpoint_dir,
79         int(best_checkpoint["epoch"]), device)
80     final_eval = evaluate(model, patch_embedding_generator,
81         test_loader, device)
82
83     save_curve(
84         training_loss_history,
85         title="CIFAR-10 visTransformer Training Loss",
86         xlabel="Epoch",
87         ylabel="Cross-Entropy Loss",
88         output_path=figures_dir / "task2_cifar_loss.png",
89         marker="o",
90     )
91     save_confusion_matrix_figure(
92         final_eval["confusion_matrix"],
93         CIFAR10_CLASSES,
94         figures_dir / "task2_cifar_confusion_matrix.png",
95         title="CIFAR-10 visTransformer Confusion Matrix",
96     )
97
98     confusion_text_lines = [
99         " " + " ".join(f"{name:>7s}" for
100             name in CIFAR10_CLASSES)]
101     for class_index, class_name in enumerate(CIFAR10_CLASSES):
102         row = final_eval["confusion_matrix"][class_index]
103         row_text = " ".join(f"{int(value):7d}" for value in row)
104         confusion_text_lines.append(f"{class_name:>8s} {row_text}")
105     (results_dir / "task2_confusion_matrix.txt").write_text("\n".join
106         (confusion_text_lines), encoding="utf-8")
107
108     summary = {
109         "config": vars(args),
110         "parameter_count": count_parameters(model) + count_parameters
111             (patch_embedding_generator),
112         "training_loss_history": training_loss_history,
113         "checkpoint_scores": checkpoint_scores,

```

Listing 2: Task 2 training, checkpointing, and evaluation wrapper around DLStudio's visTransformer

2.2 (b) Figure: Training loss curve over 40 epochs

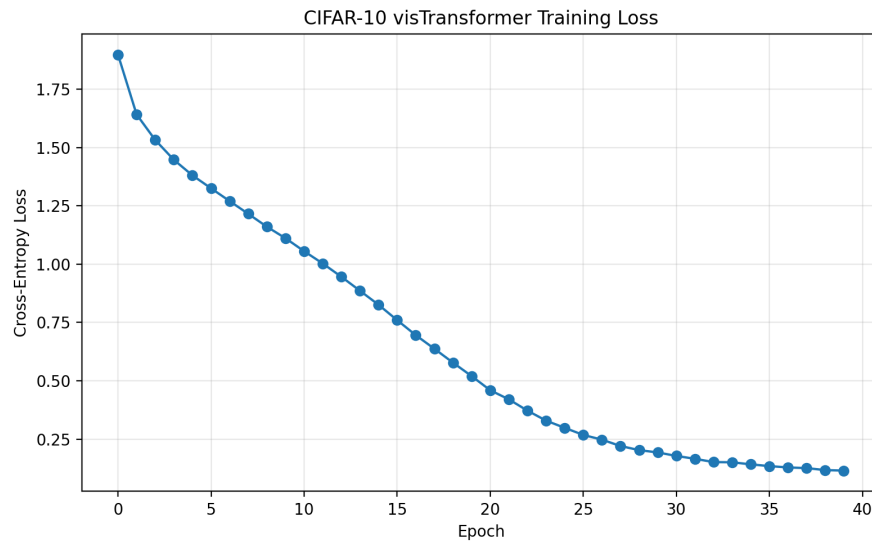


Figure 2: Task 2 CIFAR-10 visTransformer training loss over 40 epochs.

2.3 (c) Table: 10×10 confusion matrix on test set

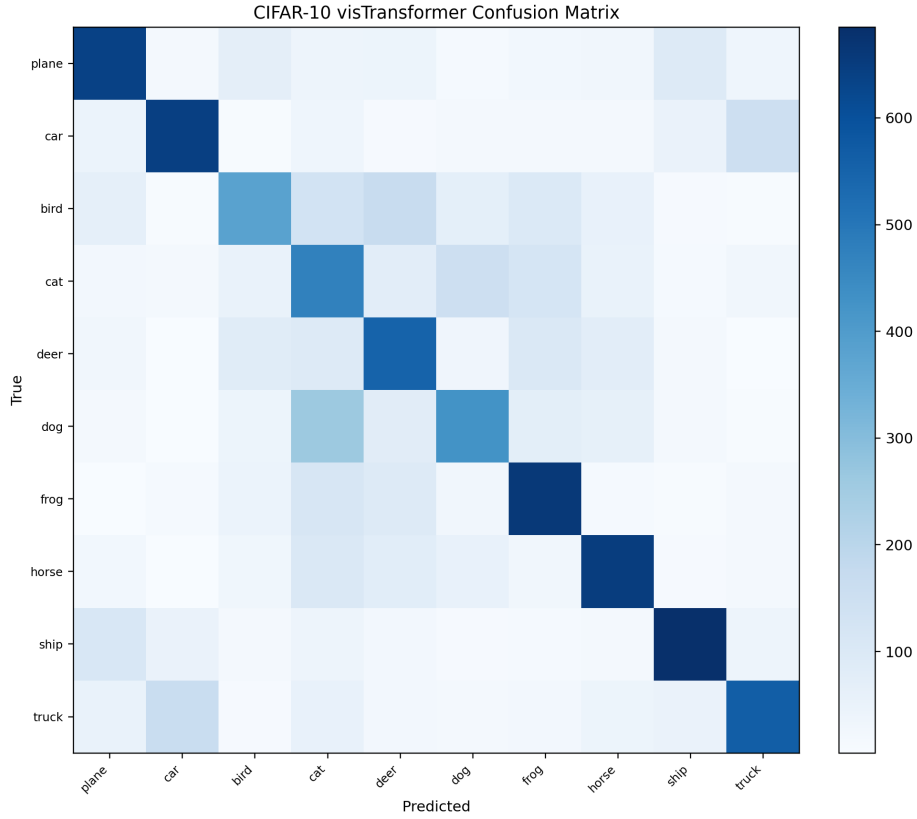


Figure 3: Task 2 confusion matrix for the best saved checkpoint (epoch 20).

Table 2: Task 2 test confusion matrix counts for the best checkpoint (epoch 20).

True \ Pred.	plane	car	bird	cat	deer	dog	frog	horse	ship	truck
plane	639	18	71	38	42	13	26	28	90	35
car	47	645	10	36	11	19	16	18	48	150
bird	66	10	381	131	167	68	99	57	12	9
cat	23	18	51	472	75	149	120	51	13	28
deer	27	5	84	94	552	34	103	76	18	7
dog	18	6	42	262	82	426	74	61	19	10
frog	5	14	46	112	95	28	659	15	10	16
horse	25	5	29	101	78	55	28	648	11	20
ship	110	49	16	39	23	11	14	16	685	37
truck	51	162	11	58	21	19	23	42	49	564

2.4 (d) Text: Overall test accuracy

Checkpoint evaluation gave the following accuracies: epoch 10: 55.47%, epoch 20: 56.71%, epoch 30: 55.79%, epoch 40: 56.24%. The best checkpoint was epoch 20, with overall test accuracy **56.71%**.

2.5 (e) Text: Patch embedding explanation

For $32 \times 32 \times 3$ CIFAR-10 images and a patch size of 16×16 , each image yields

$$P = (32/16) \times (32/16) = 4$$

patches. Each raw patch contains $16 \times 16 \times 3 = 768$ pixel values before embedding. After the linear projection to embedding width M , the patch tensor shape becomes $[4, M]$ if we ignore the batch axis.

2.6 (f) Text: Class token explanation

After concatenating one learnable class token with the 4 patch embeddings, the tensor entering the encoder has shape $[5, M]$. The first position corresponds to the class token and the remaining four correspond to image patches.

Because self-attention allows every token to exchange information with every other token at every layer, the class token acts like a learned global summary slot. By the time it emerges from the final encoder, it has repeatedly aggregated evidence from all patches and can therefore serve as a compact representation for the final classification layer.

2.7 (g) Text: Positional encoding explanation

Without positional encodings, the transformer would treat the 4 CIFAR-10 patches as an unordered set. Swapping the top-left patch with the bottom-right patch would produce the same attention computation up to permutation, even though the image semantics would change.

Positional encodings inject the notion of spatial location into the token stream. That lets the model learn distinctions such as “sky tends to be above land” or “wheels usually sit below the body of a car.” Removing positional encodings would significantly weaken the model’s ability to reason about image layout.

3 Task 3: Food-101 From Scratch + Attention Visualization (25 points)

3.1 (a) Code: Adapted visTransformer for Food-101 subset

```
1 $env:KMP_DUPLICATE_LIB_OK='TRUE'
2 conda run -n ece60146 python attention_viz.py ^
3     --mode all --epochs 20 --batch-size 64 ^
4     --results-dir results_task3 --figures-dir figures
```

Listing 3: Task 3 command used for the final from-scratch Food-20 run

```
1 class AttentionHead(nn.Module):
2     """A single self-attention head that stores its softmax weights.
3     """
4
5     def __init__(self, qkv_size: int) -> None:
6         super().__init__()
7         self.qkv_size = qkv_size
8         self.wq = nn.Linear(qkv_size, qkv_size, bias=False)
9         self.wk = nn.Linear(qkv_size, qkv_size, bias=False)
10        self.wv = nn.Linear(qkv_size, qkv_size, bias=False)
11        self.attention_weights: torch.Tensor | None = None
12
13    def forward(self, token_slice: torch.Tensor) -> torch.Tensor:
14        """Apply scaled dot-product attention to one embedding slice.
15        """
16
17        queries = self.wq(token_slice)
18        keys = self.wk(token_slice)
19        values = self.wv(token_slice)
20        scores = torch.matmul(queries, keys.transpose(1, 2)) / np.
21            sqrt(self.qkv_size)
22        weights = torch.softmax(scores, dim=-1)
23        self.attention_weights = weights.detach()
24        return torch.matmul(weights, values)
25
26
27 class SelfAttention(nn.Module):
28     """Split embeddings across heads and concatenate the attended
29     slices."""
30
31     def __init__(self, embedding_size: int, num_atten_heads: int) ->
32         None:
33         super().__init__()
34         if embedding_size % num_atten_heads != 0:
35             raise ValueError("embedding_size must be divisible by
```

```

        num_attn_heads")
31     self.embedding_size = embedding_size
32     self.num_attn_heads = num_attn_heads
33     self.qkv_size = embedding_size // num_attn_heads
34     self.attention_heads = nn.ModuleList([AttentionHead(self.
        qkv_size) for _ in range(num_attn_heads)])
35
36     def forward(self, tokens: torch.Tensor, collect_attention: bool =
        False):
37         """Return attended tokens and optionally the stacked
        attention maps."""
38
39         head_outputs = []
40         for head_index, head in enumerate(self.attention_heads):
41             token_slice = tokens[:, :, head_index * self.qkv_size : (
                head_index + 1) * self.qkv_size]
42             head_outputs.append(head(token_slice))
43         concatenated = torch.cat(head_outputs, dim=-1)
44         if not collect_attention:
45             return concatenated
46         attention_maps = torch.stack([head.attention_weights for head
            in self.attention_heads], dim=1)
47         return concatenated, attention_maps
48
49
50 class BasicEncoder(nn.Module):
51     """A residual self-attention block followed by an FFN."""
52
53     def __init__(self, embedding_size: int, num_attn_heads: int) ->
        None:
54         super().__init__()
55         self.self_attention = SelfAttention(embedding_size,
            num_attn_heads)
56         self.norm1 = nn.LayerNorm(embedding_size)
57         self.ffn = nn.Sequential(
58             nn.Linear(embedding_size, 4 * embedding_size),
59             nn.ReLU(inplace=True),
60             nn.Linear(4 * embedding_size, embedding_size),
61         )
62         self.norm2 = nn.LayerNorm(embedding_size)
63
64     def forward(self, tokens: torch.Tensor, collect_attention: bool =
        False):
65         """Forward through one encoder block."""
66
67         if collect_attention:
68             attended, attention_maps = self.self_attention(tokens,

```

```

        collect_attention=True)
69     else:
70         attended = self.self_attention(tokens, collect_attention=
            False)
71         attention_maps = None
72         tokens = self.norm1(tokens + attended)
73         tokens = self.norm2(tokens + self.ffn(tokens))
74         if collect_attention:
75             return tokens, attention_maps
76         return tokens
77
78
79 class MasterEncoder(nn.Module):
80     """Stack multiple BasicEncoder layers and optionally return
        attention maps."""
81
82     def __init__(self, num_layers: int, embedding_size: int,
        num_attn_heads: int) -> None:
83         super().__init__()
84         self.layers = nn.ModuleList(
85             [BasicEncoder(embedding_size=embedding_size,
                num_attn_heads=num_attn_heads) for _ in range(
                num_layers)]
86         )
87
88     def forward(self, tokens: torch.Tensor, collect_attention: bool =
        False):
89         """Run the full encoder stack."""
90
91         attention_weights: list[torch.Tensor] = []
92         output = tokens
93         for layer in self.layers:
94             if collect_attention:
95                 output, layer_attention = layer(output,
                    collect_attention=True)
96                 attention_weights.append(layer_attention)
97             else:
98                 output = layer(output, collect_attention=False)
99         if collect_attention:
100             return output, attention_weights
101         return output
102
103
104 class FoodVisTransformer(nn.Module):
105     """Small ViT-style classifier for the 20-class Food-101 subset.
        """
106

```

```

107     def __init__(
108         self,
109         image_size: int = 64,
110         patch_size: int = 8,
111         embedding_size: int = 128,
112         num_basic_encoders: int = 4,
113         num_attn_heads: int = 4,
114         num_classes: int = 20,
115     ) -> None:
116         super().__init__()
117         self.image_size = image_size
118         self.patch_size = patch_size
119         self.embedding_size = embedding_size
120         self.num_basic_encoders = num_basic_encoders
121         self.num_attn_heads = num_attn_heads
122         self.num_classes = num_classes
123         self.patch_embedder = PatchEmbeddingGenerator(image_size,
124                                                         patch_size, embedding_size)
125         self.class_token = nn.Parameter(torch.randn(1, 1,
126                                                         embedding_size) * 0.02)
127         self.master_encoder = MasterEncoder(num_basic_encoders,
128                                                         embedding_size, num_attn_heads)
129         self.classifier = nn.Sequential(
130             nn.LayerNorm(embedding_size),
131             nn.Linear(embedding_size, 256),
132             nn.GELU(),
133             nn.Dropout(0.1),
134             nn.Linear(256, num_classes),
135         )
136
137     def prepare_tokens(self, images: torch.Tensor) -> torch.Tensor:
138         """Build '[cls] + patch_embeddings' tokens from raw images.
139         """
140
141         patch_embeddings = self.patch_embedder(images)
142         class_tokens = self.class_token.expand(images.shape[0], -1,
143                                                 -1)
144         return torch.cat((class_tokens, patch_embeddings), dim=1)
145
146     def forward_tokens(self, tokens: torch.Tensor, collect_attention:
147                        bool = False):
148         """Run the encoder stack on prebuilt tokens."""
149
150         if collect_attention:
151             encoded, attention_weights = self.master_encoder(tokens,
152                                                             collect_attention=True)
153             logits = self.classifier(encoded[:, 0])

```

```

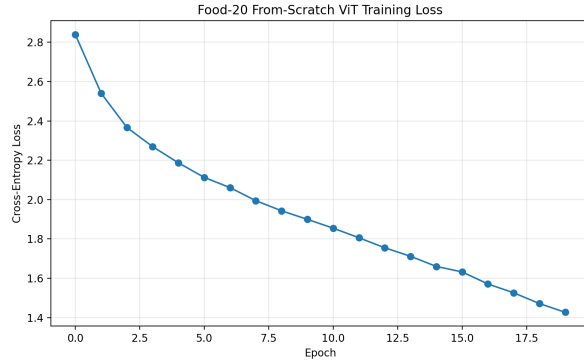
147         return logits, attention_weights
148     encoded = self.master_encoder(tokens, collect_attention=False)
149     return self.classifier(encoded[:, 0])
150
151     def forward(self, x: torch.Tensor, collect_attention: bool =
152     False):
153         """Accept either images '[B, 3, H, W]' or token sequences '[B
154         , S, M]'."""
155
156         if x.dim() == 4:
157             tokens = self.prepare_tokens(x)
158         elif x.dim() == 3:
159             tokens = x
160         else:
161             raise ValueError("Expected images with 4 dims or token
162             sequences with 3 dims.")
163         return self.forward_tokens(tokens, collect_attention=
164         collect_attention)
165
166 def extract_attention_weights(model: FoodVisTransformer, input_tensor
167 : torch.Tensor) -> list[torch.Tensor]:
168     """Return one '[heads, seq_len, seq_len]' tensor per encoder
169     layer.
170
171     The autograder supplies a token tensor shaped like '[1, seq_len,
172     embedding_size]',
173     so this function uses the encoder path directly and returns the
174     first batch item.
175     """
176
177     device = next(model.parameters()).device
178     model.eval()
179     with torch.no_grad():
180         _, attention_weights = model(input_tensor.to(device),
181         collect_attention=True)
182     return [layer_weights[0].detach().cpu() for layer_weights in
183     attention_weights]

```

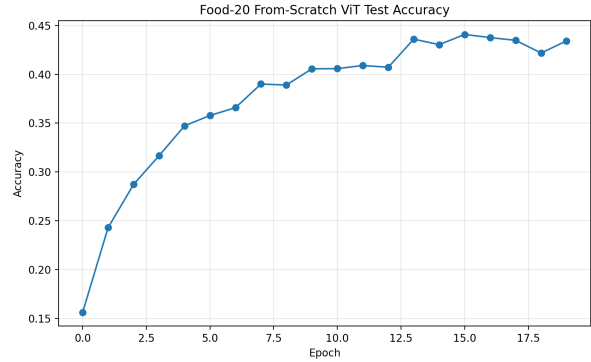
Listing 4: Modified attention stack, encoder, model, and required autograder hook in `attention_viz.py`

The adapted model uses the exact 20-class Food-101 subset from the assignment, resizes images to 64×64 , and uses 8×8 patches. Therefore each image is converted to 64 patch tokens, and after prepending the class token the sequence length is 65, which produces a 65×65 attention map per head. The final classifier outputs 20 logits. The final from-scratch model had 649,492 learnable parameters.

3.2 (b) Figure: Training loss curve



(a) Training loss



(b) Test accuracy per epoch

Figure 4: Task 3 from-scratch Food-20 training curves.

3.3 (c) Text: Test accuracy and training observations

The model converged steadily from 15.62% test accuracy after epoch 1 to a best of **44.08%** at epoch 16. Since random guessing over 20 classes would only give 5%, the trained model is clearly learning meaningful category structure even though it is far below the pretrained ViT in Task 4.

3.4 (d) Code: Modified AttentionHead, SelfAttention, BasicEncoder, and MasterEncoder

```
1 class AttentionHead(nn.Module):
2     """A single self-attention head that stores its softmax weights.
3     """
4
5     def __init__(self, qkv_size: int) -> None:
6         super().__init__()
7         self.qkv_size = qkv_size
8         self.wq = nn.Linear(qkv_size, qkv_size, bias=False)
9         self.wk = nn.Linear(qkv_size, qkv_size, bias=False)
10        self.wv = nn.Linear(qkv_size, qkv_size, bias=False)
11        self.attention_weights: torch.Tensor | None = None
12
13    def forward(self, token_slice: torch.Tensor) -> torch.Tensor:
14        """Apply scaled dot-product attention to one embedding slice.
15        """
16
17        queries = self.wq(token_slice)
18        keys = self.wk(token_slice)
19        values = self.wv(token_slice)
```

```

18         scores = torch.matmul(queries, keys.transpose(1, 2)) / np.
           sqrt(self.qkv_size)
19         weights = torch.softmax(scores, dim=-1)
20         self.attention_weights = weights.detach()
21         return torch.matmul(weights, values)
22
23
24 class SelfAttention(nn.Module):
25     """Split embeddings across heads and concatenate the attended
       slices."""
26
27     def __init__(self, embedding_size: int, num_attn_heads: int) ->
       None:
28         super().__init__()
29         if embedding_size % num_attn_heads != 0:
30             raise ValueError("embedding_size must be divisible by
               num_attn_heads")
31         self.embedding_size = embedding_size
32         self.num_attn_heads = num_attn_heads
33         self.qkv_size = embedding_size // num_attn_heads
34         self.attention_heads = nn.ModuleList([AttentionHead(self.
               qkv_size) for _ in range(num_attn_heads)])
35
36     def forward(self, tokens: torch.Tensor, collect_attention: bool =
       False):
37         """Return attended tokens and optionally the stacked
           attention maps."""
38
39         head_outputs = []
40         for head_index, head in enumerate(self.attention_heads):
41             token_slice = tokens[:, :, head_index * self.qkv_size : (
               head_index + 1) * self.qkv_size]
42             head_outputs.append(head(token_slice))
43         concatenated = torch.cat(head_outputs, dim=-1)
44         if not collect_attention:
45             return concatenated
46         attention_maps = torch.stack([head.attention_weights for head
               in self.attention_heads], dim=1)
47         return concatenated, attention_maps
48
49
50 class BasicEncoder(nn.Module):
51     """A residual self-attention block followed by an FFN."""
52
53     def __init__(self, embedding_size: int, num_attn_heads: int) ->
       None:
54         super().__init__()

```

```

55     self.self_attention = SelfAttention(embedding_size,
56                                         num_attn_heads)
57     self.norm1 = nn.LayerNorm(embedding_size)
58     self.ffn = nn.Sequential(
59         nn.Linear(embedding_size, 4 * embedding_size),
60         nn.ReLU(inplace=True),
61         nn.Linear(4 * embedding_size, embedding_size),
62     )
63     self.norm2 = nn.LayerNorm(embedding_size)
64
65     def forward(self, tokens: torch.Tensor, collect_attention: bool =
66                 False):
67         """Forward through one encoder block."""
68
69         if collect_attention:
70             attended, attention_maps = self.self_attention(tokens,
71                                                             collect_attention=True)
72         else:
73             attended = self.self_attention(tokens, collect_attention=
74                                             False)
75             attention_maps = None
76         tokens = self.norm1(tokens + attended)
77         tokens = self.norm2(tokens + self.ffn(tokens))
78         if collect_attention:
79             return tokens, attention_maps
80         return tokens
81
82     class MasterEncoder(nn.Module):
83         """Stack multiple BasicEncoder layers and optionally return
84         attention maps."""
85
86         def __init__(self, num_layers: int, embedding_size: int,
87                     num_attn_heads: int) -> None:
88             super().__init__()
89             self.layers = nn.ModuleList(
90                 [BasicEncoder(embedding_size=embedding_size,
91                               num_attn_heads=num_attn_heads) for _ in range(
92                     num_layers)]
93             )
94
95         def forward(self, tokens: torch.Tensor, collect_attention: bool =
96                     False):
97             """Run the full encoder stack."""
98
99             attention_weights: list[torch.Tensor] = []
100             output = tokens

```

```

93         for layer in self.layers:
94             if collect_attention:
95                 output, layer_attention = layer(output,
96                     collect_attention=True)
97                 attention_weights.append(layer_attention)
98             else:
99                 output = layer(output, collect_attention=False)
100         if collect_attention:
101             return output, attention_weights
102         return output
103
104 class FoodVisTransformer(nn.Module):
105     """Small ViT-style classifier for the 20-class Food-101 subset.
106     """

```

Listing 5: AttentionHead with stored weights, SelfAttention, BasicEncoder, and MasterEncoder with optional attention collection

3.5 (e) Text: Description of attention extraction modifications

I added an `attention_weights` instance attribute to `AttentionHead` immediately after the row-wise softmax. Then `SelfAttention` stacks those stored maps into a tensor of shape `[batch, heads, seq, seq]`, and `BasicEncoder` / `MasterEncoder` thread those tensors upward only when `collect_attention=True`. The top-level helper `extract_attention_weights(model, input_tensor)` runs a no-grad forward pass and returns a list of `[heads, seq, seq]` tensors, one per encoder layer, exactly as required by the autograder.

These changes do not interfere with training because they only detach and read the softmax weights. The classifier path remains unchanged: the class token still summarizes the encoded sequence and the final linear head still predicts one of the 20 food categories.

3.6 (f) Figure: Part A — 4 images with full heatmap and spatial overlay

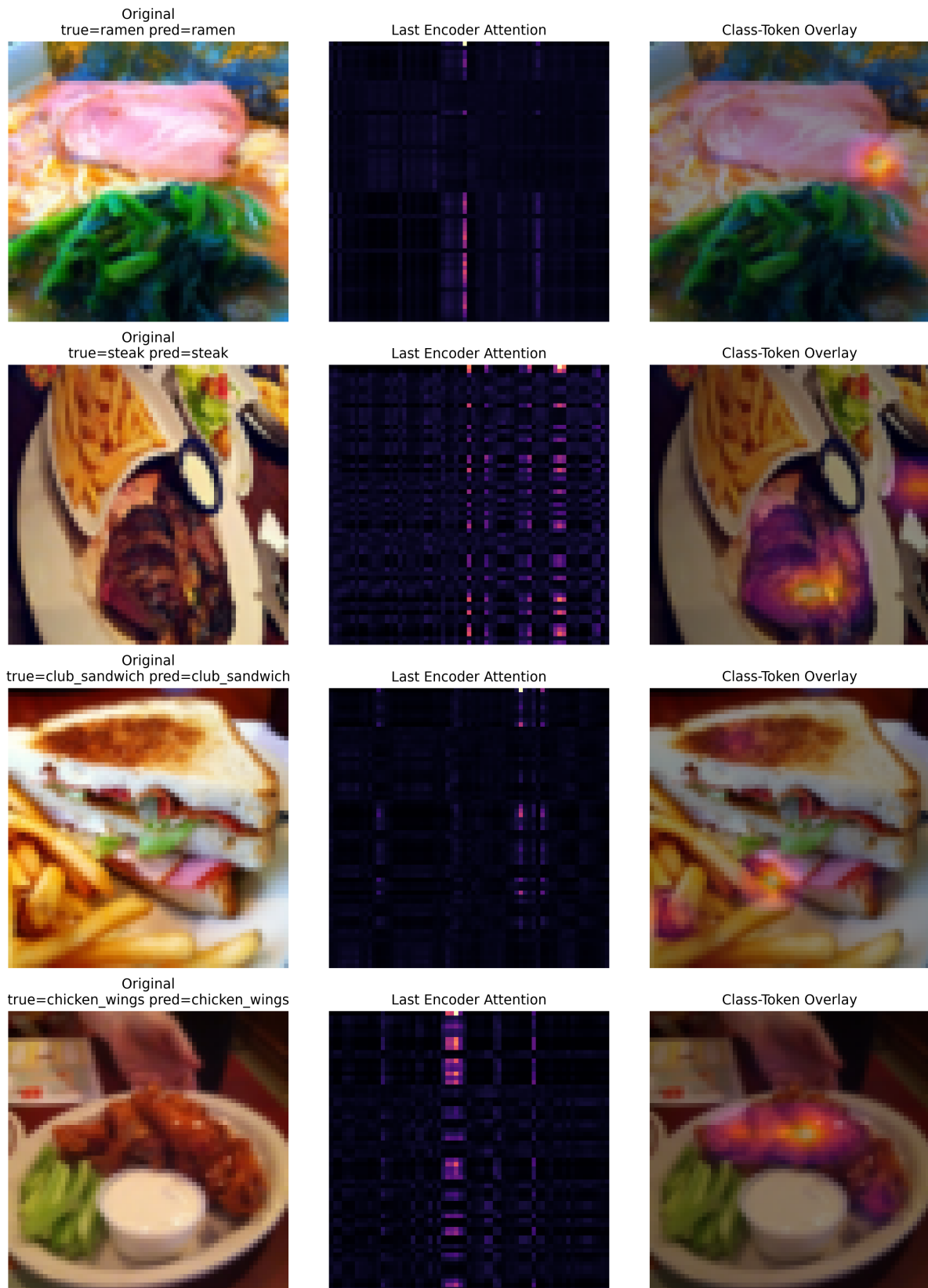


Figure 5: Task 3 Part A. For each correctly classified image: original image, full 65×65 last-layer attention map for one head, and class-token spatial overlay.

3.7 (g) Text: Part A — Per-image attention analysis

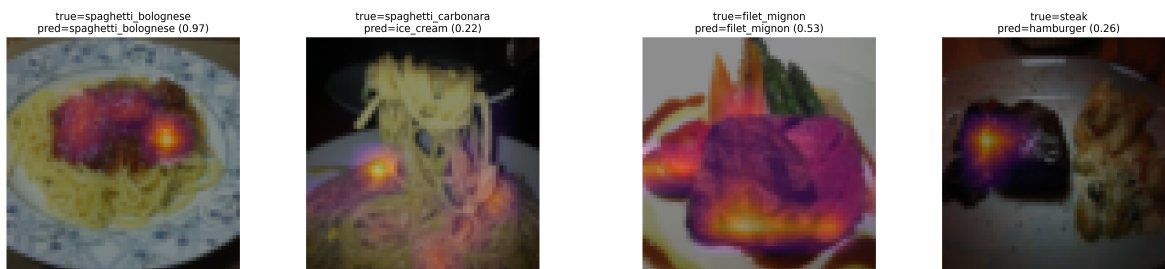
Ramen. The overlay concentrates on the bright garnish and pork region on the right side of the bowl while still covering part of the noodles and broth. That suggests the class token is using a visually distinctive topping pattern rather than uniformly attending to the entire bowl.

Steak. The class token focuses on the browned meat surface in the lower part of the plate and mostly ignores the fries. That is a sensible decision boundary because the seared meat texture is more class-specific than the side dish.

Club sandwich. The strongest attention falls near the sandwich interior and plate garnish region, where the layered bread, filling, and color contrast are most obvious. The fries occupy a large area in the image, but the overlay treats them as secondary context.

Chicken wings. The map emphasizes the glossy wing surface across the upper-middle and right side of the plate while largely ignoring the dipping sauce cup and lettuce garnish. The model appears to be using the sauce-coated wing texture and curved wing shapes as the main evidence.

3.8 (h) Figure: Part B — 2 pairs with side-by-side overlays



(a) Spaghetti bolognese vs. spaghetti carbonara

(b) Filet mignon vs. steak

Figure 6: Task 3 Part B side-by-side class-token attention overlays for two visually similar category pairs.

3.9 (i) Text: Part B — Comparative analysis per pair

Spaghetti bolognese vs. spaghetti carbonara. For the bolognese image, the class token concentrates on the red meat-sauce mound at the center-right, which is exactly the feature that differentiates it from plain or cream-based pasta dishes. For the carbonara image, the attention centers on the pale creamy region and glossy highlights rather than on clear pasta structure, and the model misclassifies it as `ice_cream`. That failure suggests the low-color-contrast, smooth cream texture may be confusing the model more than the actual noodle layout. The two overlays confirm the model is responding to sauce color and texture rather than to the shape of the noodles themselves.

Filet mignon vs. steak. The filet mignon example receives broad attention across the meat surface, especially the lower seared edge and central body of the cut, and it is classified correctly. In contrast, the steak example is misclassified as **hamburger**, and the overlay locks onto a dark rectangular slab of meat while paying relatively little attention to the surrounding potatoes. That makes sense: the attended region looks somewhat patty-like, so the model appears to be collapsing the two red-meat categories when the cut shape is ambiguous. A clearer plating or a different camera angle might have produced the correct prediction.

4 Task 4: Pretrained ViT Fine-tuning on Food-101 Subset (45 points)

4.1 (a) Code: Fine-tuning script for ViT-Base on Food-101 subset

```
1 def build_loaders(batch_size: int, data_root: str | Path) -> tuple[
2     DataLoader, DataLoader]:
3     """Create train and test loaders with ImageNet-style transforms.
4     """
5
6     train_transform = T.Compose(
7         [
8             T.RandomResizedCrop(224),
9             T.RandomHorizontalFlip(),
10            T.ToTensor(),
11            T.Normalize(IMAGENET_MEAN, IMAGENET_STD),
12        ]
13    )
14    test_transform = T.Compose(
15        [
16            T.Resize(256),
17            T.CenterCrop(224),
18            T.ToTensor(),
19            T.Normalize(IMAGENET_MEAN, IMAGENET_STD),
20        ]
21    )
22    train_bundle = make_food20("train", train_transform, root=
23        data_root)
24    test_bundle = make_food20("test", test_transform, root=data_root)
25    train_loader = DataLoader(train_bundle.dataset, batch_size=
26        batch_size, shuffle=True, num_workers=4, pin_memory=True)
27    test_loader = DataLoader(test_bundle.dataset, batch_size=
28        batch_size, shuffle=False, num_workers=4, pin_memory=True)
29    return train_loader, test_loader
30
31 def maybe_freeze_backbone(model: nn.Module, freeze: bool) -> None:
32     """Optionally freeze the patch embed and transformer blocks."""
33
34     for parameter in model.patch_embed.parameters():
35         parameter.requires_grad = not freeze
36     for block in model.blocks:
37         for parameter in block.parameters():
38             parameter.requires_grad = not freeze
```

```

37 def evaluate(model: nn.Module, loader: DataLoader, device: torch.
    device) -> dict[str, Any]:
38     """Evaluate a classifier and keep 16 examples for the prediction
        grid."""
39
40     model.eval()
41     correct = 0
42     total = 0
43     dataset_size = len(loader.dataset)
44     target_positions = set(np.linspace(0, dataset_size - 1, 16, dtype
        =int).tolist())
45     global_index = 0
46     sampled_images: list[torch.Tensor] = []
47     sampled_predictions: list[str] = []
48     sampled_truths: list[str] = []
49     with torch.no_grad():
50         for images, labels in loader:
51             images = images.to(device, non_blocking=True)
52             labels = labels.to(device, non_blocking=True)
53             logits = model(images)
54             predictions = torch.argmax(logits, dim=1)
55             correct += (predictions == labels).sum().item()
56             total += labels.numel()
57             for image, prediction, truth in zip(images.cpu(),
                predictions.cpu(), labels.cpu()):
58                 if global_index in target_positions:
59                     sampled_images.append(image)
60                     sampled_predictions.append(FOOD20_CLASSES[int(
                        prediction)])
61                     sampled_truths.append(FOOD20_CLASSES[int(truth)])
62             global_index += 1
63     return {
64         "accuracy": correct / max(total, 1),
65         "sampled_images": sampled_images,
66         "sampled_predictions": sampled_predictions,
67         "sampled_truths": sampled_truths,
68     }

```

Listing 6: Task 4 data transforms, evaluation helper, and 16-image grid sampling logic

```

1 def main() -> None:
2     """Fine-tune the pretrained ViT-Base model and save all report
        artifacts."""
3
4     args = parse_args()
5     setup_seed(args.seed)
6     results_dir = ensure_dir(PROJECT_ROOT / args.results_dir)
7     figures_dir = ensure_dir(PROJECT_ROOT / args.figures_dir)

```

```

8     checkpoint_path = results_dir / "finetuned_vit_best.pt"
9
10    device = torch.device("cuda" if torch.cuda.is_available() else "
        cpu")
11    train_loader, test_loader = build_loaders(args.batch_size,
        PROJECT_ROOT / args.data_root)
12    model = timm.create_model("vit_base_patch16_224", pretrained=True
        , num_classes=len(FOOD20_CLASSES)).to(device)
13
14    criterion = nn.CrossEntropyLoss()
15    train_loss_history: list[float] = []
16    validation_accuracy_history: list[float] = []
17    best_accuracy = -1.0
18    start_time = time.perf_counter()
19
20    for epoch in range(1, args.epochs + 1):
21        maybe_freeze_backbone(model, freeze=epoch <= args.
            freeze_epochs)
22        optimizer = torch.optim.AdamW(
23            [parameter for parameter in model.parameters() if
                parameter.requires_grad],
24            lr=args.learning_rate,
25            weight_decay=args.weight_decay,
26        )
27
28        model.train()
29        running_loss = 0.0
30        example_count = 0
31        for images, labels in train_loader:
32            images = images.to(device, non_blocking=True)
33            labels = labels.to(device, non_blocking=True)
34            optimizer.zero_grad()
35            logits = model(images)
36            loss = criterion(logits, labels)
37            loss.backward()
38            optimizer.step()
39            running_loss += loss.item() * labels.shape[0]
40            example_count += labels.shape[0]
41
42        mean_loss = running_loss / max(example_count, 1)
43        validation = evaluate(model, test_loader, device)
44        validation_accuracy = validation["accuracy"]
45        train_loss_history.append(mean_loss)
46        validation_accuracy_history.append(validation_accuracy)
47        elapsed = time.perf_counter() - start_time
48        print(
49            f"epoch {epoch:02d}/{args.epochs} train_loss={mean_loss

```

```

50         :.4f} "
        f"val_accuracy={validation_accuracy:.4%} elapsed={elapsed
51         :.1f}s"
52     )
53     if validation_accuracy > best_accuracy:
54         best_accuracy = validation_accuracy
55         torch.save(
56             {
57                 "model_state": model.state_dict(),
58                 "config": vars(args),
59                 "best_accuracy": best_accuracy,
60             },
61             checkpoint_path,
62         )
63
64     checkpoint = torch.load(checkpoint_path, map_location=device)
65     model.load_state_dict(checkpoint["model_state"])
66     model.to(device)
67     final_eval = evaluate(model, test_loader, device)
68
69     save_dual_curve(
70         train_loss_history,
71         validation_accuracy_history,
72         first_label="Training Loss",
73         second_label="Validation Accuracy",
74         title="ViT-Base Fine-Tuning Curves",
75         xlabel="Epoch",
76         ylabel="Value",
77         output_path=figures_dir / "task4_vit_curves.png",
78     )
79     save_prediction_grid(
80         final_eval["sampled_images"],
81         final_eval["sampled_predictions"],
82         final_eval["sampled_truths"],
83         figures_dir / "task4_vit_prediction_grid.png",
84         mean=IMAGENET_MEAN,
85         std=IMAGENET_STD,
86     )
87
88     summary = {
89         "config": vars(args),
90         "train_loss_history": train_loss_history,
91         "validation_accuracy_history": validation_accuracy_history,
92         "best_accuracy": best_accuracy,
93         "test_accuracy": final_eval["accuracy"],
94         "parameter_count": count_parameters(model),

```

4.2 (b) Figure: Training loss and validation accuracy curves

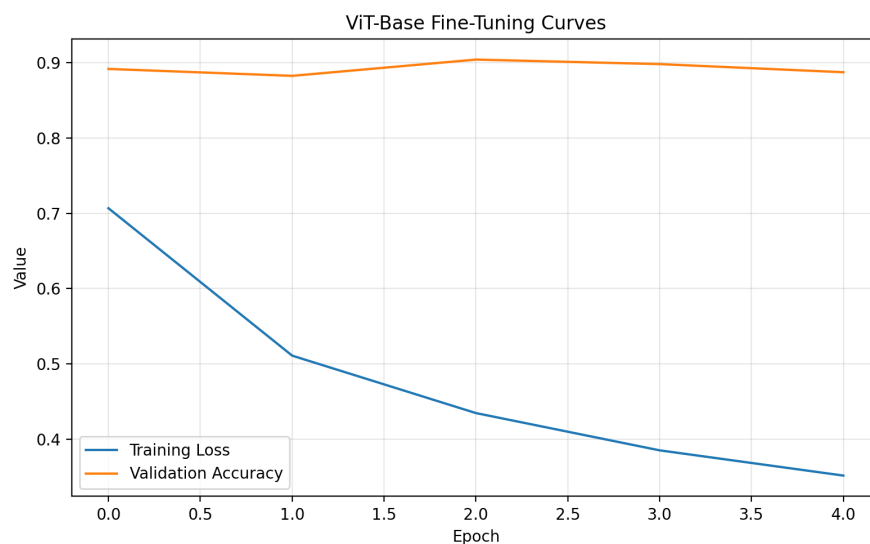


Figure 7: Task 4 training loss and validation accuracy curves over 5 fine-tuning epochs.

4.3 (c) Text: Test accuracy and parameter count

The best fine-tuned ViT-Base checkpoint reached **90.40%** accuracy on the Food-20 test split, with the peak occurring at epoch 3. The model has **85,814,036** learnable parameters.

4.4 (d) Figure: 4×4 image grid with predictions

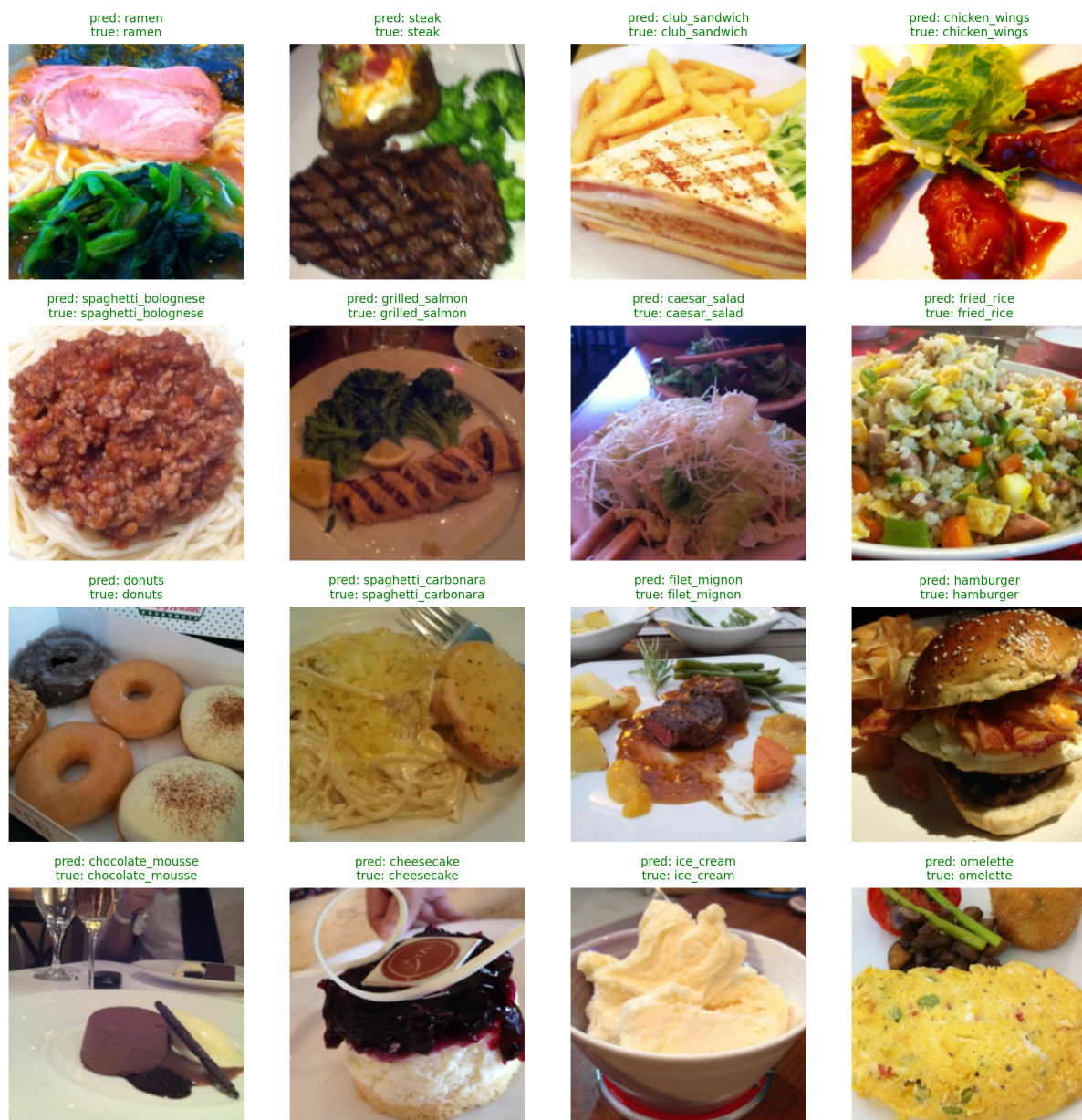


Figure 8: Task 4 4×4 test-image grid with predicted and true labels. Incorrect predictions are shown in red; the saved sample here is entirely correct.

4.5 (e) Table: Two-model comparison table

Table 3: Comparison of the from-scratch Food-20 model and the pretrained ViT-Base.

Model	Input Size	Pretrained?	Accuracy	Params
visTransformer from Task 3	64 ²	No	44.08%	649,492
ViT-Base fine-tuned from Task 4	224 ²	Yes	90.40%	85,814,036

4.6 (f) Text: Accuracy gap discussion

The accuracy gap is $90.40\% - 44.08\% = 46.32$ percentage points. All three listed factors matter: the pretrained model is larger, it sees higher-resolution inputs, and it starts from pretrained weights. However, I believe pretraining contributes the most. The from-scratch Food-20 transformer already reaches 44% with only about 0.65M parameters and 64×64 images, which means the task is learnable from the subset. The dramatic jump to 90% is hard to explain by model size or input resolution alone; the pretrained weights give the model a mature feature hierarchy and much better token-to-token relational priors before fine-tuning even begins.

4.7 (g) Text: Model capacity and data requirements discussion

ViT-Base has roughly 85.8M trainable parameters in this run, compared with about 0.65M for the Task 3 model. If I trained ViT-Base from scratch on only 15,000 Food-20 training images, I would expect severe overfitting and optimization difficulty. Unlike a CNN, a ViT has weak built-in inductive bias for locality, so a large randomly initialized transformer usually needs massive pretraining data to learn stable patch representations and inter-patch relationships. In other words, a huge uninitialized ViT would likely do worse than the much smaller Task 3 model because the data budget is too small for its capacity.

4.8 (h) Text: Connection to language model pretraining

This homework made the pretraining idea very concrete. In vision, ImageNet pretraining teaches the transformer reusable structure such as edges, textures, part-whole relationships, and category-level semantics. Fine-tuning on Food-20 then only needs to specialize those already-useful features to the target classes.

The language-model analogy is direct. BERT and GPT are pretrained on very large generic text corpora so they learn syntax, word meaning, long-range context, and token prediction structure before a downstream task is introduced. Fine-tuning on a smaller labeled dataset works well because the model is not learning language from scratch anymore; it is just adapting a rich prior to a narrower objective. ViT on ImageNet and language models on large corpora are the same story in two different modalities.