



INTERNATIONAL INSTITUTE OF
INFORMATION TECHNOLOGY

HYDERABAD



BharatGen

GenAI for Bharat, by Bharat

<https://bharatgen.com>

Making Vision Models Answer and Reason

Ravi Kiran Sarvadevabhatla
CVIT, IIIT Hyderabad



@vikataravi

Evolution of Vision-Language Models

Vision piggybacked on advances in language generation.

Evolution
in NLP

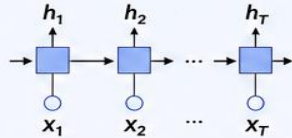
1990s–
2000s

1

Advances in Language Generation (NLP)

LSTM

Model long
sequences



CNN
(Visual Encoder)



LSTM
(Language Model)



Caption

"A dog is
running."

Image Captioning

How Vision Used It (Vision-Language Models)

Adopted
by Vision

 Vision learns
to describe
images

Evolution of Vision-Language Models

Vision piggybacked on advances in language generation.

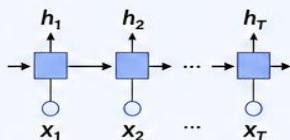
Evolution
in NLP

1990s–
2000s

1

LSTM

Model long
sequences



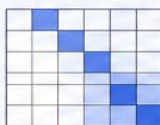
2010s
early

2

Attention

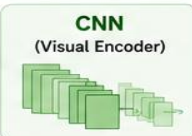
Focus on the
most relevant
parts

Attention weights



Advances in Language Generation (NLP)

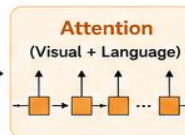
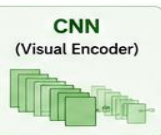
How Vision Used It (Vision-Language Models)



Caption

"A dog is running."

Image Captioning



Caption

"A dog is running."

Show, Attend and Tell

Adopted
by Vision



Vision learns
to describe
images



Vision learns
to focus on
important
regions

Evolution of Vision-Language Models

Vision piggybacked on advances in language generation.

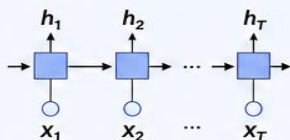
Evolution in NLP

1990s–
2000s

1

LSTM

Model long sequences



2010s
early

2

Attention

Focus on the most relevant parts

Attention weights

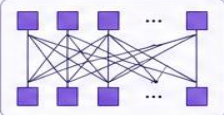


2010s
mid–late

3

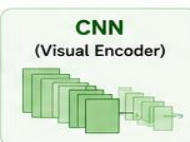
Transformer

Model global dependencies with self-attention



Advances in Language Generation (NLP)

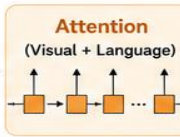
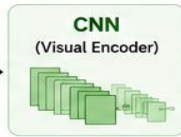
How Vision Used It (Vision-Language Models)



Caption

“A dog is running.”

Image Captioning



Caption

“A dog is running.”

Show, Attend and Tell



Text

“A dog is running.”

Vision Transformers + Transformer Decoder

Adopted by Vision



Vision learns to describe images



Vision learns to focus on important regions



Vision + Transformers model complex visual structures

Evolution of Vision-Language Models

Vision piggybacked on advances in language generation.

Evolution in NLP

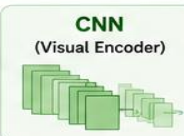
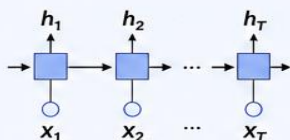
Advances in Language Generation (NLP)

How Vision Used It (Vision-Language Models)

Adopted by Vision

1 LSTM

Model long sequences



LSTM (Language Model)

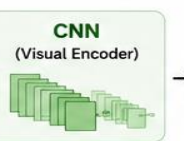
Caption

"A dog is running."

Image Captioning

2 Attention

Focus on the most relevant parts



Attention (Visual + Language)

LSTM (Language Model)

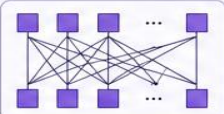
Caption

"A dog is running."

Show, Attend and Tell

3 Transformer

Model global dependencies with self-attention



Transformer Decoder

Text

"A dog is running."

Vision Transformers + Transformer Decoder

4 LLM

Generate text, reason, and converse



LLM (Large Language Model)

Conversation / Answer

Q: What is happening here?
A: A person is boating on a lake surrounded by mountains.

Vision Encoder + LLM (Modern VLMs)

Vision learns to describe images

Vision learns to focus on important regions

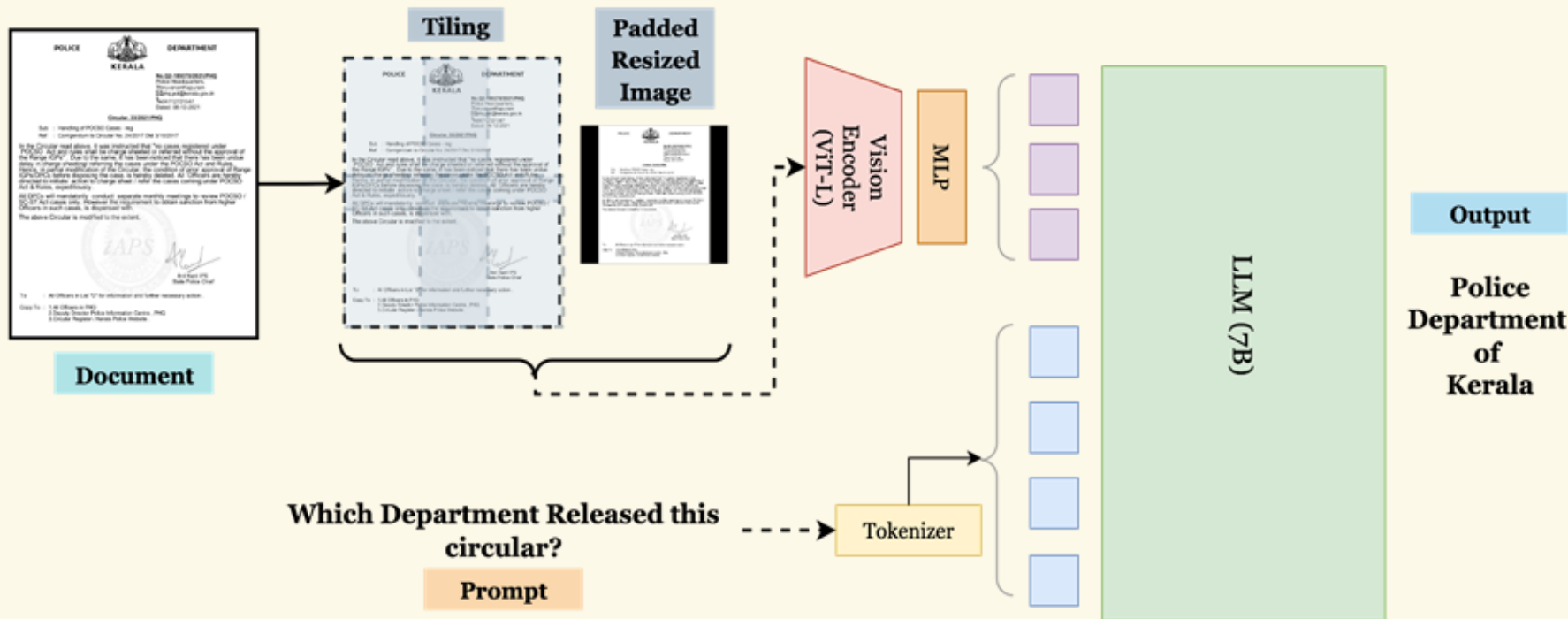
Vision + Transformers model complex visual structures

Vision + LLMs enable rich conversations and reasoning



Each generation of Vision-Language models **inherited the latest advances** in language generation.

Document Visual Question-Answering (DocVQA)



Patram - India's First Document VLM

What is the IFSC Code in this Cheque

Indian Bank
Branch: TALOJA
SHOP: 12 AND 3 AMAR HARMONY PLOT NO 22 SECTOR 4
TALOJA PHASE 1 NEAR PANCHAND
RAJWADI STATION NAVI MUMBAI 410208
IFSC Code: IDIB000T199

PAY *Self Rohan Kumar* OR BEARER
या धारक को

RUPEES रुपये *पाँच हजार रुपये मात्र*
अदा करें ₹ *5000/-*

CBS Code: 3028

Text: *Dependra A. Jha*
Please sign above

980700030

PAYABLE AT PAR AT ALL OUR BRANCHES

BB2018 4000190831 287490* 31



PATRAM



Gemma 3

Correct Answer

IDIB000T199

Hallucination

3028

Document VLM Leaderboard: Patram Results

Model	Overall	DocVQA	Patram-Bench	VisualMRC
Qwen2.5-VL-7B-Instruct	0.8759	0.8722	0.6816	0.9169
Patram-7B-Instruct	0.8692	0.9373	0.6888	0.8598
Gemma-3-12B-IT	0.8556	0.8451	0.6349	0.9069
InternVL3-9B	0.7865	0.8681	0.6432	0.7405
DeepSeek-VL2	0.7581	0.8739	0.5089	0.7144
Molmo-7B-D-0924	0.6820	0.6314	0.5018	0.7577
DocOwl2	0.6776	0.7391	0.4569	0.5868
Molmo-7B-O-0924	0.6654	0.6314	0.5018	0.7577
Llava-v1.6-mistral-7b	0.2577	0.1263	0.4134	0.3889



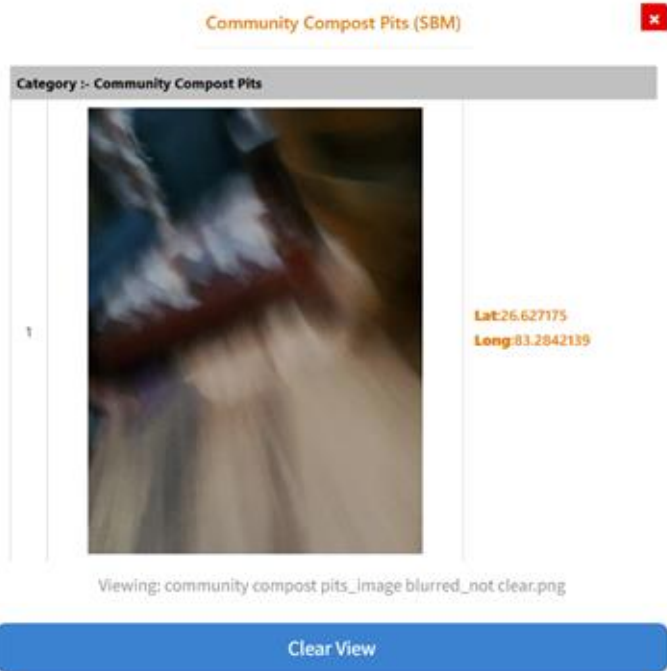
DocEval
An Evaluation framework for DocAI Tasks



<https://huggingface.co/bharatgenai/patram-7b-instruct>

WASH Usecase: Compost Pit Pictures

Example 1: Image Blur + Lat Long extraction



Chat with Doc



Namaste Bharatgen! Welcome to your new DocBodh chat. How can I assist you?



is the image in the document blurred or clear?



The image in the document is blurred and abstract, making it difficult to identify specific details or elements.



can you give the lat and long?

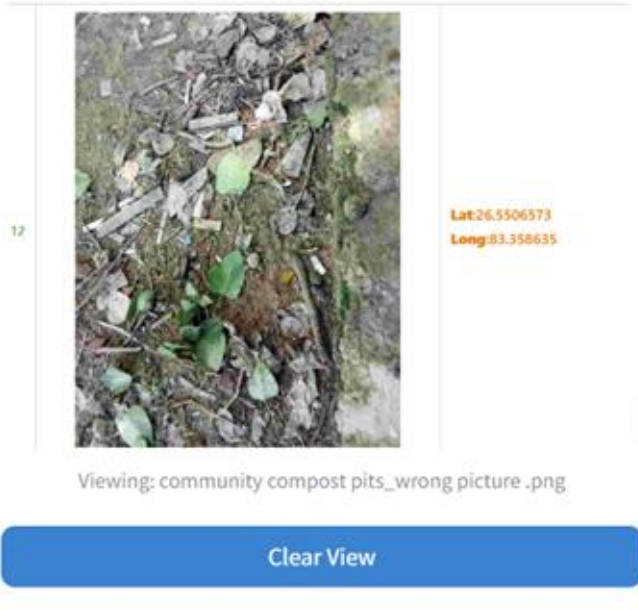
[Attached:  community compost pits_image blurred_not clear.png]



The latitude is 26.627175 and the longitude is 83.2842139.

WASH Usecase: Compost Pit Pictures

Example 2 - Image is compost pit or not + Lat long extraction



Chat with Doc



Namaste Bharatgen! Welcome to your new DocBodh chat. How can I assist you?



is this a picture of a compost pit inside the document?



The image does not show a compost pit. It depicts a natural, cluttered area with debris and vegetation.



what is the lat and long?



The latitude is 26.5506573 and the longitude is 83.358635.

In a Document VQA system, *Answers Must Be Verifiable!*

Q What can be inferred from the low values in this blood report?

A The report shows a low **MEAN CORPUSCULAR VOLUME, MCV** of **72.0 fL** which indicates **Microcytic anemia**



? *"Which line in this report does that come from? Can you show me?"*

HAEMATOLOGY			
COMPLETE BLOOD COUNT (CBC)			
TEST	VALUE	UNIT	REFERENCE
HEMOGLOBIN	15	g/dl	13 - 17
TOTAL LEUKOCYTE COUNT	5,500	count	4,800 - 10,800
DIFFERENTIAL LEUCOCYTE COUNT			
NEUTROPHILS	79	%	40 - 80
LYMPHOCYTE	18	%	20 - 40
EOSINOPHILS	1	%	1 - 6
MONOCYTES	4	%	2 - 10
BASOPHILS	1	%	0 - 2
PLATELET COUNT	3.5	lakhs/cumm	1.5 - 4.5
TOTAL RBC COUNT	5	million/cumm	4.5 - 5.5
HEMATOCRIT VALUE, HCT	42	%	40 - 50
MEAN CORPUSCULAR VOLUME, MCV	L 72.0	fL	83 - 101
MEAN CELL HAEMOGLOBIN, MCH	30.0	Pg	27 - 32
MEAN CELL HAEMOGLOBIN CON, MCHC	32.7	%	31.5 - 34.5

Clinical Notes:
A complete blood count (CBC) is used to evaluate overall health and detect a wide range of disorders, including anemia, infection, and leukimia. There have been some reports of RBC and platelet counts being lower in venous blood than in capillary blood samples, although still within these reference ranges.

Possible causes of abnormal parameters:

High	Low	
RBC, Hb, or HCT	Dehydration, polycythemia, shock, chronic hypoxia	Anemia, thalassemia, and other hemoglobinopathies
MCV	Macrocytic anemia, liver disease	Microcytic anemia
WBC (Leucocytes)	Acute stress, infection, malignancies	Infection risk, weak immune system
Platelets	Risk of thrombosis	Risk of bleeding

Mr. Sachin Sharma
DMLT, Lab Incharge

Dr. A. K. Anshu
MBBS, MD Pathologist

Page 1 of 2

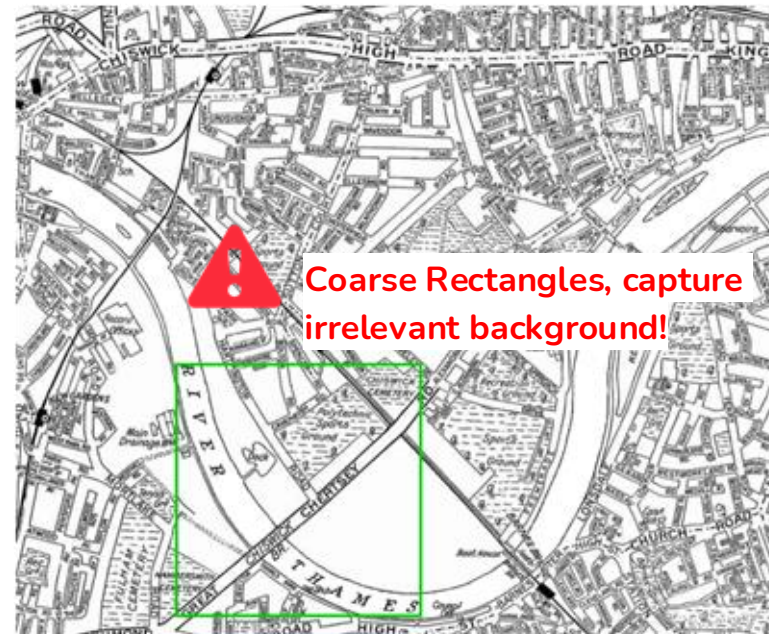
NOT VALID FOR MEDICO-LEGAL PURPOSE
Work hours: Monday to Sunday, 9 am to 6 pm

Please consult clinically. Although the test results are checked thoroughly, in case of any unexpected test results which could be due to machine error or typing error or any other reason please contact the lab immediately for a free evaluation.



Problem with Current DocVQA Approaches

1. Do not ground.
2. Box-based grounding is coarse and fails on curved/complex text.
3. Prior approaches lack multi-span and multi-granular grounding.



What is the name of the river in the map?



River Thames [x1,y1,x2,y2]



What M3Grounder Does

Add a **smart highlighter** to your VLM !

A

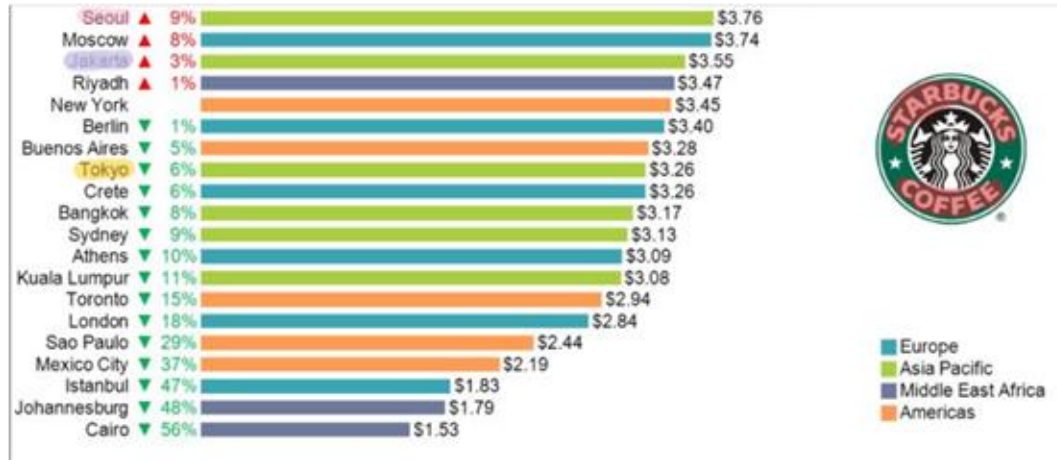


What is the name of the river in the map?



`<e>River Thames</e>[GROUND].`

B



According to the chart, which coffee brand is being compared, and which three Asia Pacific cities have the highest average coffee prices?



The coffee is `<e>Starbucks Coffee</e>[GROUND]`, the cities with the highest prices in Asia Pacific are `<e>Seoul</e>[GROUND]`, `<e>Jakarta</e>[GROUND]`, `<e>Tokyo</e>[GROUND]`.



What M3Grounder Does

C1

Phrase Level

TCF NATIONAL BANK
1410 KENTON LN W
PLANO, TX 75041

STATEMENT

STATEMENT DATE
08-24-17
8812847839

STATEMENT PERIOD
07-27-17 THROUGH 08-31-17

ACCOUNT NUMBER
8812847839

YOU HAVE OPTED-OUT OF TCP'S AUTORIZATION AND PAYMENT OF OVERDRAFTS ON YOUR ATM AND EVERYDAY DEBIT CARD TRANSACTIONS. YOU HAVE OPTED-OUT OF TCP'S PAYMENT OF OVERDRAFTS ON YOUR CHECKS, ELECTRONIC TRANSACTIONS, AND TRANSFERS. SEE THE REVERSE SIDE FOR MORE INFORMATION.

ACCOUNT SUMMARY BALANCE 07-24-17 08-31-17

DATE	AMOUNT	DESCRIPTION
08-01	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-02	210.00	POS FRANCISCA'S 8832 - Clothing Store
08-03	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-04	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-05	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-06	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-07	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-08	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-09	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-10	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-11	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-12	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-13	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-14	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-15	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-16	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-17	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-18	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-19	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-20	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-21	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-22	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-23	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-24	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-25	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-26	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-27	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-28	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-29	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-30	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-31	120.00	POS FRANCISCA'S 8832 - Clothing Store

FOR BALANCE AND CHECKS PAID INFORMATION, VISIT TCP WEBSITE, VISIT US ONLINE AT TCP.NA.COM OR CALL 1-800-431-6141. (ISSUING INSTRUCTIONS) YOU CAN ALSO DIRECT INQUIRIES TO THE ADDRESS SHOWN AT THE TOP OF THIS PAGE. FOR ALL ACCOUNTS OTHER THAN TCP CHECKS (CHECKING), TCP CHARGES \$17 FOR OVERDRAFTS AND RETURNED ITEMS. FOR TCP CHECKS (CHECKING), TCP CHARGES \$5 FOR EACH DAY YOUR ACCOUNT IS OVERDRAWN.

C2

Line Level

TCF NATIONAL BANK
1410 KENTON LN W
PLANO, TX 75041

STATEMENT

STATEMENT DATE
08-24-17
8812847839

STATEMENT PERIOD
07-27-17 THROUGH 08-31-17

ACCOUNT NUMBER
8812847839

YOU HAVE OPTED-OUT OF TCP'S AUTORIZATION AND PAYMENT OF OVERDRAFTS ON YOUR ATM AND EVERYDAY DEBIT CARD TRANSACTIONS. YOU HAVE OPTED-OUT OF TCP'S PAYMENT OF OVERDRAFTS ON YOUR CHECKS, ELECTRONIC TRANSACTIONS, AND TRANSFERS. SEE THE REVERSE SIDE FOR MORE INFORMATION.

ACCOUNT SUMMARY BALANCE 07-24-17 08-31-17

DATE	AMOUNT	DESCRIPTION
08-01	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-02	210.00	POS FRANCISCA'S 8832 - Clothing Store
08-03	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-04	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-05	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-06	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-07	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-08	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-09	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-10	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-11	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-12	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-13	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-14	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-15	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-16	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-17	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-18	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-19	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-20	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-21	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-22	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-23	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-24	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-25	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-26	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-27	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-28	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-29	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-30	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-31	120.00	POS FRANCISCA'S 8832 - Clothing Store

FOR BALANCE AND CHECKS PAID INFORMATION, VISIT TCP WEBSITE, VISIT US ONLINE AT TCP.NA.COM OR CALL 1-800-431-6141. (ISSUING INSTRUCTIONS) YOU CAN ALSO DIRECT INQUIRIES TO THE ADDRESS SHOWN AT THE TOP OF THIS PAGE. FOR ALL ACCOUNTS OTHER THAN TCP CHECKS (CHECKING), TCP CHARGES \$17 FOR OVERDRAFTS AND RETURNED ITEMS. FOR TCP CHECKS (CHECKING), TCP CHARGES \$5 FOR EACH DAY YOUR ACCOUNT IS OVERDRAWN.

C3

Block Level

TCF NATIONAL BANK
1410 KENTON LN W
PLANO, TX 75041

STATEMENT

STATEMENT DATE
08-24-17
8812847839

STATEMENT PERIOD
07-27-17 THROUGH 08-31-17

ACCOUNT NUMBER
8812847839

YOU HAVE OPTED-OUT OF TCP'S AUTORIZATION AND PAYMENT OF OVERDRAFTS ON YOUR ATM AND EVERYDAY DEBIT CARD TRANSACTIONS. YOU HAVE OPTED-OUT OF TCP'S PAYMENT OF OVERDRAFTS ON YOUR CHECKS, ELECTRONIC TRANSACTIONS, AND TRANSFERS. SEE THE REVERSE SIDE FOR MORE INFORMATION.

ACCOUNT SUMMARY BALANCE 07-24-17 08-31-17

DATE	AMOUNT	DESCRIPTION
08-01	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-02	210.00	POS FRANCISCA'S 8832 - Clothing Store
08-03	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-04	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-05	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-06	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-07	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-08	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-09	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-10	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-11	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-12	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-13	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-14	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-15	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-16	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-17	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-18	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-19	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-20	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-21	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-22	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-23	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-24	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-25	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-26	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-27	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-28	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-29	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-30	120.00	POS FRANCISCA'S 8832 - Clothing Store
08-31	120.00	POS FRANCISCA'S 8832 - Clothing Store

FOR BALANCE AND CHECKS PAID INFORMATION, VISIT TCP WEBSITE, VISIT US ONLINE AT TCP.NA.COM OR CALL 1-800-431-6141. (ISSUING INSTRUCTIONS) YOU CAN ALSO DIRECT INQUIRIES TO THE ADDRESS SHOWN AT THE TOP OF THIS PAGE. FOR ALL ACCOUNTS OTHER THAN TCP CHECKS (CHECKING), TCP CHARGES \$17 FOR OVERDRAFTS AND RETURNED ITEMS. FOR TCP CHECKS (CHECKING), TCP CHARGES \$5 FOR EACH DAY YOUR ACCOUNT IS OVERDRAWN.



What are my top 3 highest spends this month?



Your three highest spends are $\leq \$210.00$, $\leq \$120.00$, and $\leq \$120.00$.



Answering and Locating, Simultaneously



Question q

What can be inferred from the low values in this blood report?
Provide the answer with grounding



Answer w/ Interleaved Grounding Tokens $\hat{a}$

The report shows a low **<e>MEAN CORPUSCULAR VOLUME, MCV of 72.0 fL </e>[GROUND]** which indicates **<e>Microcytic anemia</e> [GROUND]**

Mr. Saubik Bhattacharya		Registered as: 171010004 28 04 19	Consent as: 171010004
Age/Sex: 27M		Received as: 171010004	Received as: 171010004
Ref No: 12345		Received as: 171010004	Received as: 171010004

HAEMATOLOGY			
COMPLETE BLOOD COUNT (CBC)			
TEST	VALUE	UNIT	REFERENCE
HEMOGLOBIN	15	g/dL	13 - 17
TOTAL LEUKOCYTE COUNT	5,100	cumm	4,800 - 10,800
DIFFERENTIAL LEUCOCYTE COUNT			
NEUTROPHILS	79	%	40 - 80
LYMPHOCYTE	20	%	20 - 40
EOSINOPHILS	1	%	1 - 6
MONOCYTES	4	%	2 - 10
BASOPHILS	1	%	< 2
PLATELET COUNT	3.5	lakhs/cumm	1.5 - 4.1
TOTAL RBC COUNT	5	million/cumm	4.5 - 5.5
HEMATOCRIT VALUE, HCT	42	%	40 - 50
MEAN CORPUSCULAR VOLUME, MCV	84.0	fL	83 - 101
MEAN CELL HAEMOGLOBIN, MCH	30.0	Pg	27 - 32
MEAN CELL HAEMOGLOBIN CON, MCHC	35.7	%	31.5 - 34.5

Clinical Notes:
A complete blood count (CBC) is used to evaluate overall health and detect a wide range of disorders, including anemia, infection, and leukemia. There have been some reports of MCH and platelet counts being lower in various blood than in capillary blood samples, although still within these reference ranges.

Possible causes of abnormal parameters:

High	Low
RBC, Hb, or HCT: Dehydration, polycythemia, shock, chronic hypoxia	Anemia, thalassemia, and other hemoglobinopathies
MCV: Macrocytic anemia, liver disease	Microcytic anemia
WBC (Leucocytes): Acute stress, infection, malignancies	Infection risk, weak immune system
Platelets: Risk of thrombosis	Risk of bleeding

Dr. S. S. Bhattacharya
MD, MB, BSc (Hons)

Document Image x

TEST	VALUE	UNIT	REFERENCE
HEMOGLOBIN	15	g/dL	13 - 17
TOTAL LEUKOCYTE COUNT	5,100	cumm	4,800 - 10,800
DIFFERENTIAL LEUCOCYTE COUNT			
NEUTROPHILS	79	%	40 - 80
LYMPHOCYTE	20	%	20 - 40
EOSINOPHILS	1	%	1 - 6
MONOCYTES	4	%	2 - 10
BASOPHILS	1	%	< 2
PLATELET COUNT	3.5	lakhs/cumm	1.5 - 4.1
TOTAL RBC COUNT	5	million/cumm	4.5 - 5.5
HEMATOCRIT VALUE, HCT	42	%	40 - 50
MEAN CORPUSCULAR VOLUME, MCV	84.0	fL	83 - 101
MEAN CELL HAEMOGLOBIN, MCH	30.0	Pg	27 - 32
MEAN CELL HAEMOGLOBIN CON, MCHC	35.7	%	31.5 - 34.5

Clinical Notes:
A complete blood count (CBC) is used to evaluate overall health and detect a wide range of disorders, including anemia, infection, and leukemia. There have been some reports of MCH and platelet counts being lower in various blood than in capillary blood samples, although still within these reference ranges.

Possible causes of abnormal parameters:

High	Low
RBC, Hb, or HCT: Dehydration, polycythemia, shock, chronic hypoxia	Anemia, thalassemia, and other hemoglobinopathies
MCV: Macrocytic anemia, liver disease	Microcytic anemia
WBC (Leucocytes): Acute stress, infection, malignancies	Infection risk, weak immune system
Platelets: Risk of thrombosis	Risk of bleeding

Phrase, Line, Block

Phrase, Line, Block



Final Multi-Span & Multi Granular Masks $\hat{M}$
(Overlaid on Image)

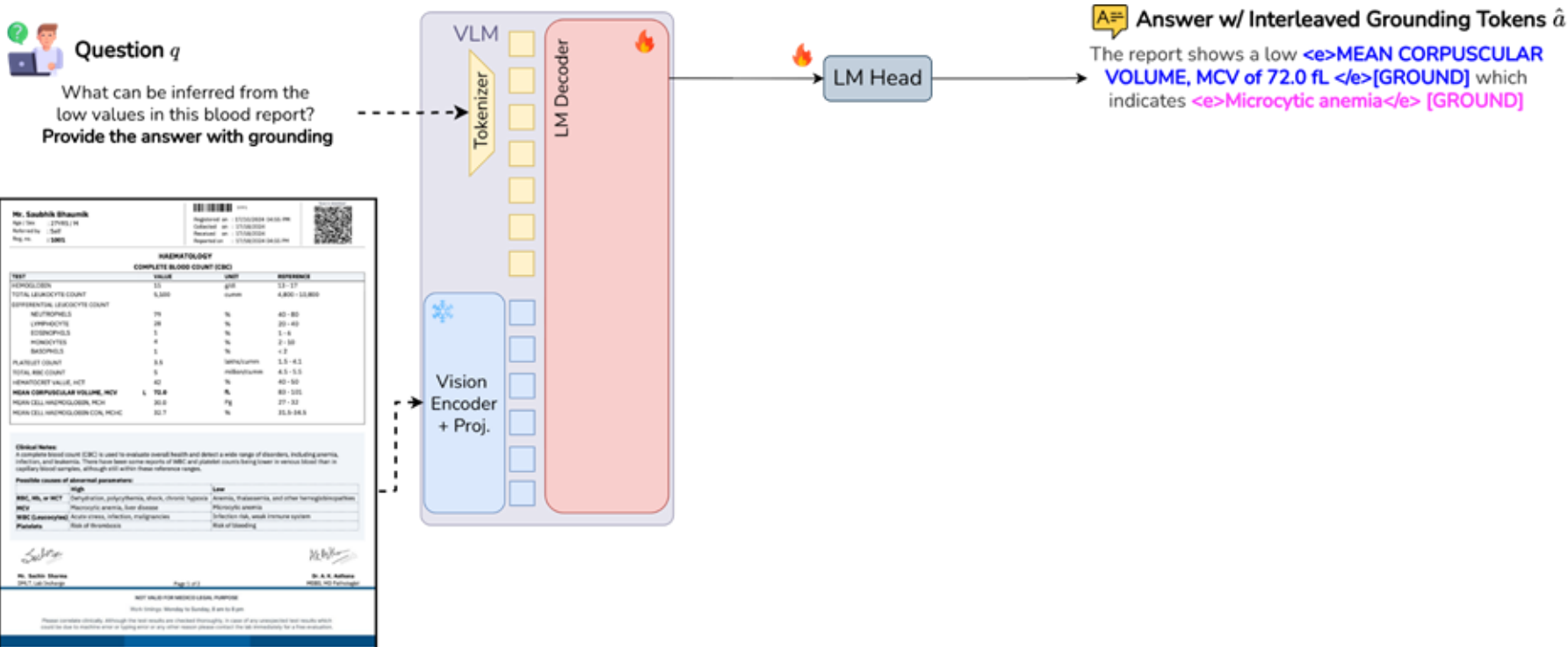


What can be inferred from the low values in this blood report?
Provide the answer with grounding

Document Image x



M3Grounder Architecture



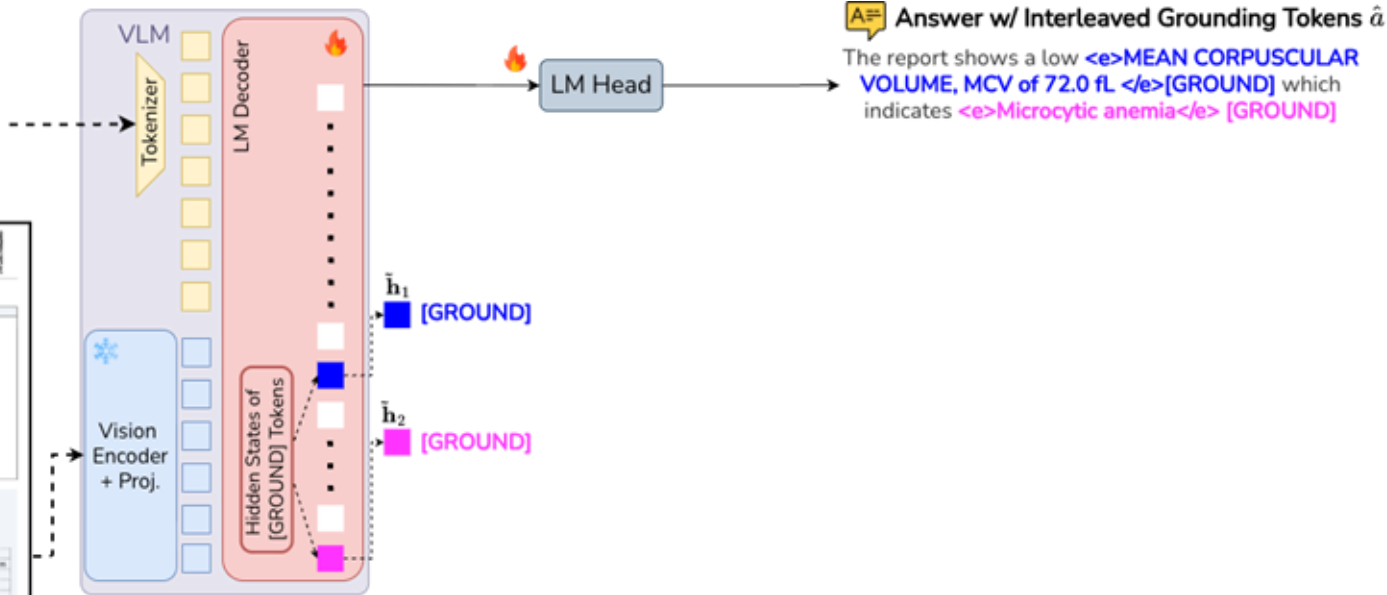


Question q

What can be inferred from the low values in this blood report?
Provide the answer with grounding

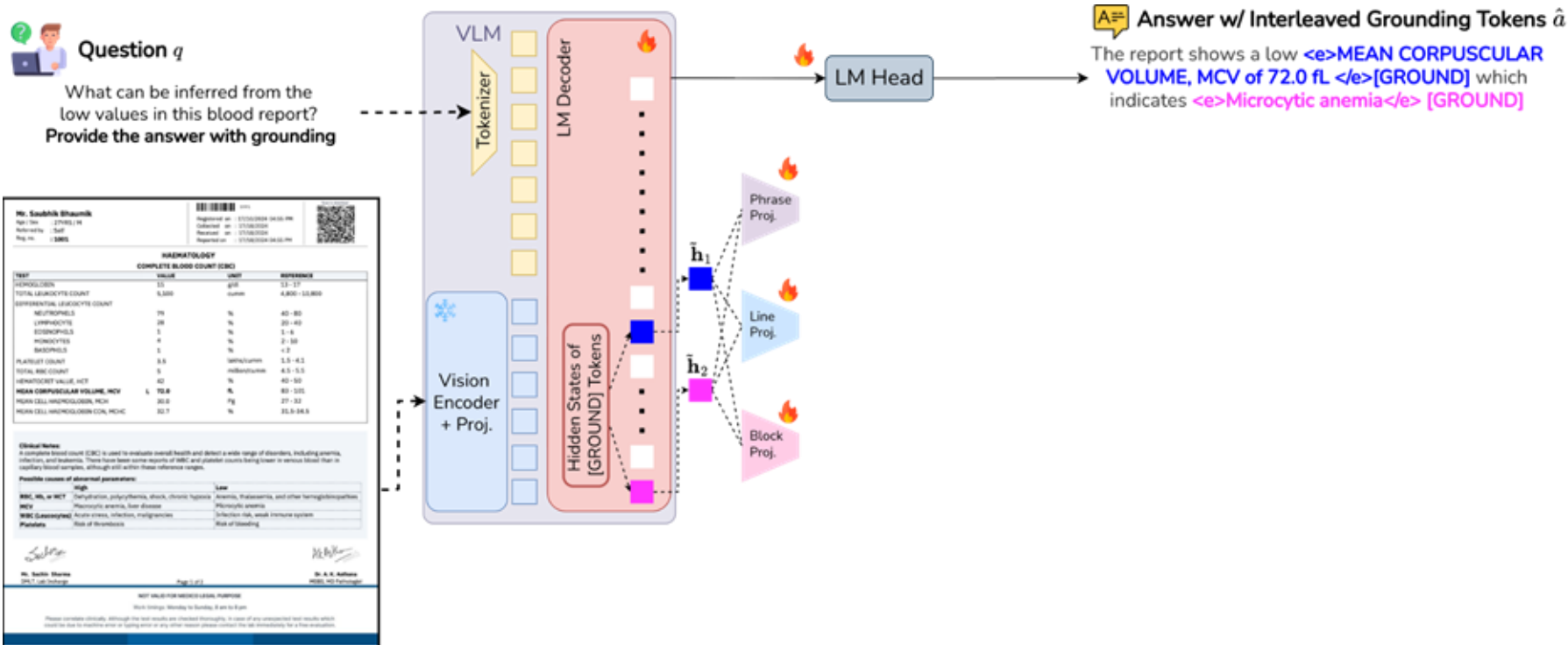
[illegible]

Document Image x



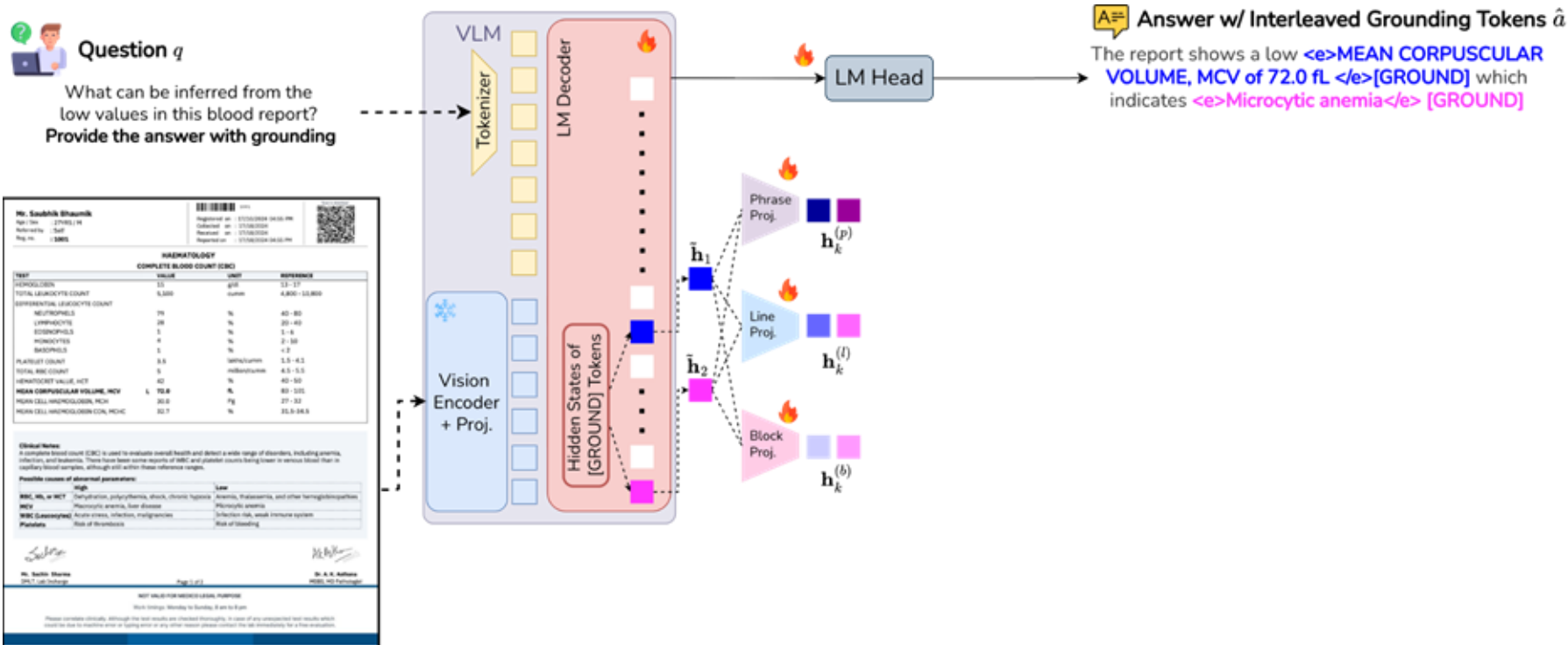


M3Grounder Architecture





M3Grounder Architecture





M3Grounder Architecture



Question q

What can be inferred from the low values in this blood report?
Provide the answer with grounding

Mr. Saubhik Bhattacharya
Reg. No. : 27001, 19
Referral No. : 5407
Reg. No. : 3005

Signature on : 17/10/2024 04:02 PM
Checked on : 17/10/2024
Reviewed on : 17/10/2024
Reported on : 17/10/2024 04:02 PM

HEMATOLOGY
COMPLETE BLOOD COUNT (CBC)

TEST	VALUE	UNIT	REFERENCE
HEMATOCRIT	42	g/dl	3.7 - 4.7
TOTAL LEUCOCYTES COUNT	5,500	count	4,800 - 10,800
DIFFERENTIAL LEUCOCYTES COUNT			
NEUTROPHILS	79	%	40 - 80
LYMPHOCYTES	20	%	20 - 40
EOSINOPHILS	1	%	1 - 6
MONOCYTES	4	%	2 - 10
PLATELETS	1	%	1 - 2
PLATELET COUNT	0.5	mill/cmm	1.5 - 4.5
TOTAL RBC COUNT	5	mill/cmm	4.5 - 5.5
HEMATOCRIT VALUE, HCT	42	%	40 - 50
MEAN CORPUSCULAR VOLUME, MCV	82.8	fL	80 - 100
MEAN CELL HEMOGLOBIN, MCH	30.0	Pg	27 - 32
MEAN CELL HEMOGLOBIN CON, MCHC	32.7	%	31.9 - 34.9

Clinical Note:

A complete blood count (CBC) is used to evaluate overall health and detect a wide range of disorders, including anemia, infection, and leukemia. These have been some reports of MCH and platelet counts being lower in venous blood than in capillary blood samples, although still within these reference ranges.

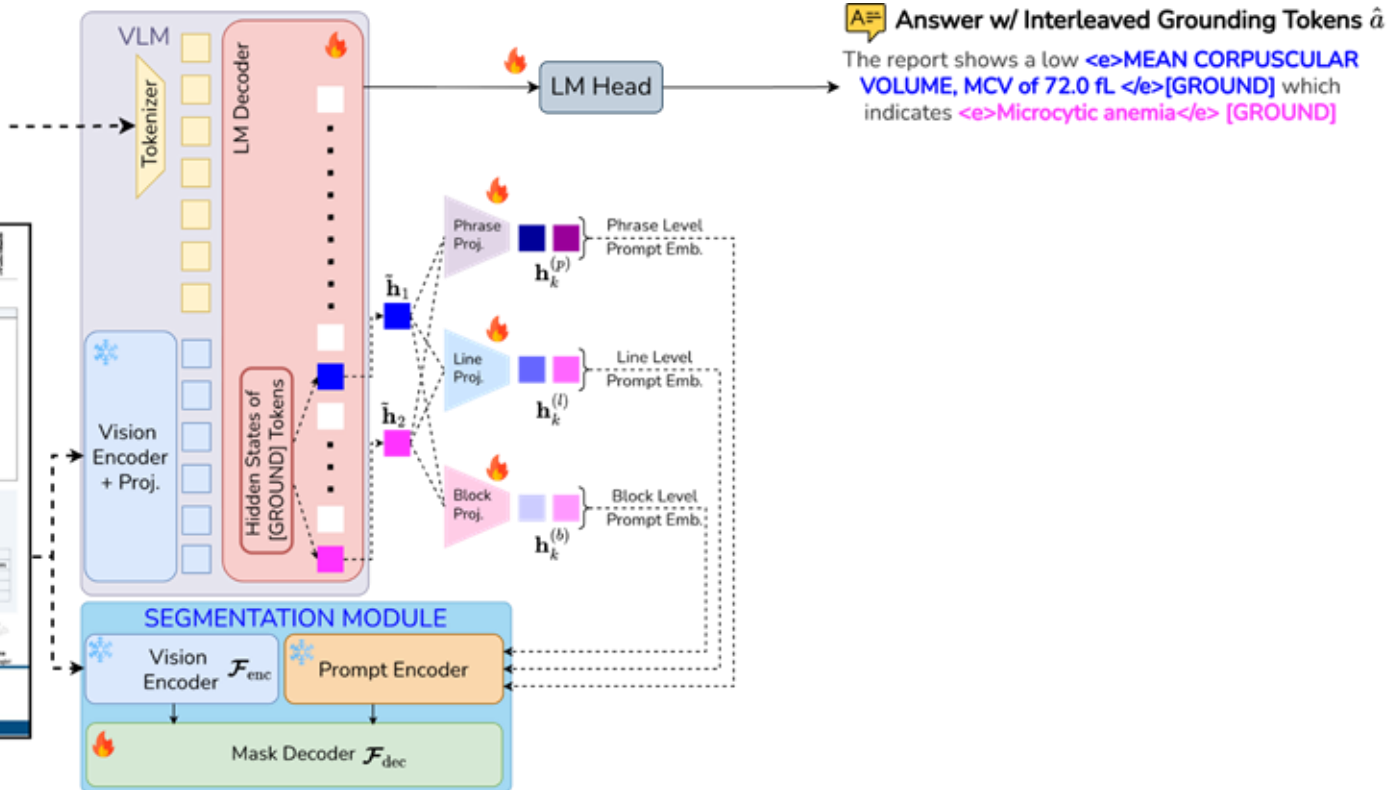
Possible causes of abnormal parameters:

High	Low	
MCH, MCV, HCT	Dehydration, polycythemia, shock, chronic hypoxia	Anemia, thalassemia, and other hemoglobinopathies
MCH	Microcytic anemia, liver disease	Microcytic anemia
MCHC (corrected)	Acute stress, infection, malignancy	Infection risk, small intestine system
Platelets	Risk of thrombosis	Risk of bleeding

Dr. Saubhik Bhattacharya
Reg. No. : 27001, 19
Referral No. : 5407
Reg. No. : 3005

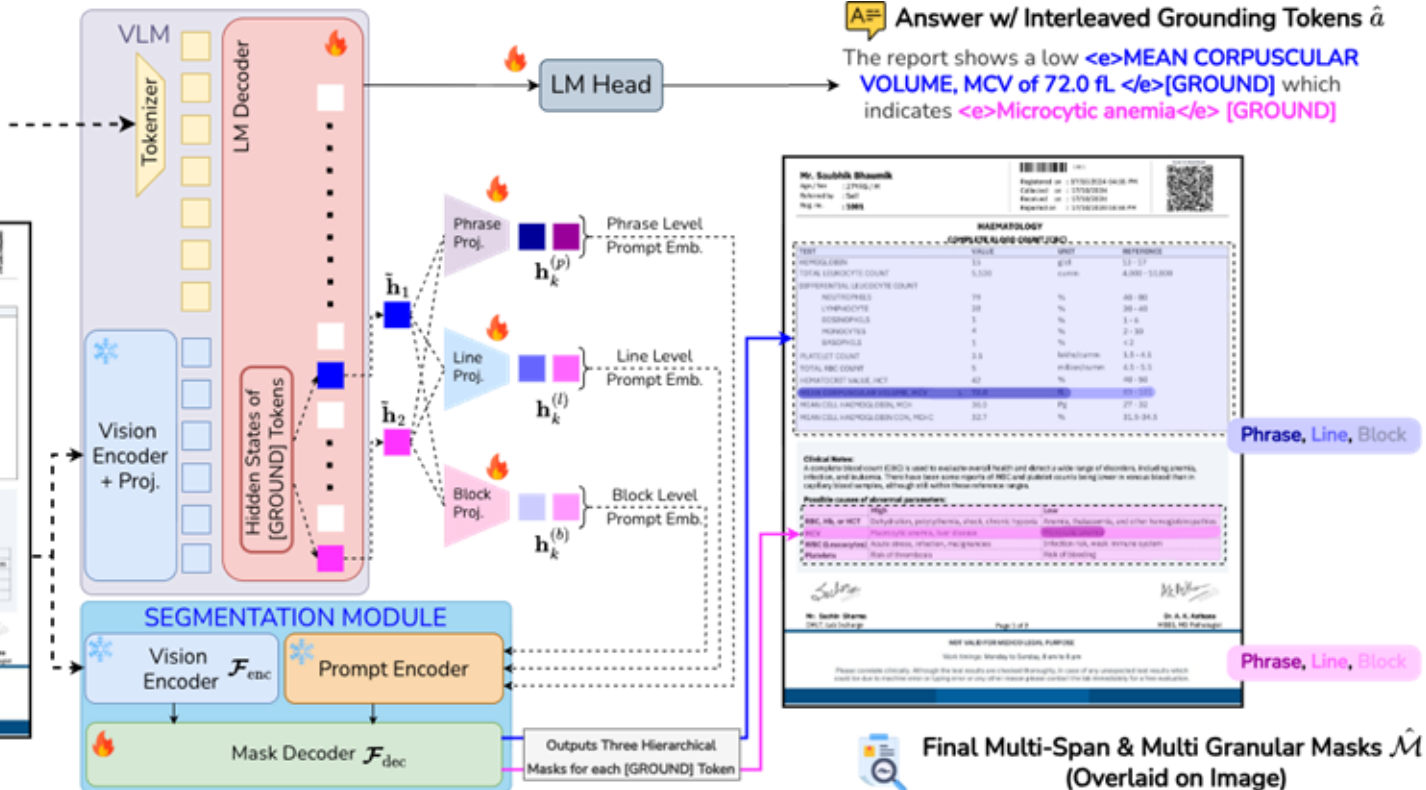
Dr. A. A. Ashraf
Reg. No. : 000

Document Image x



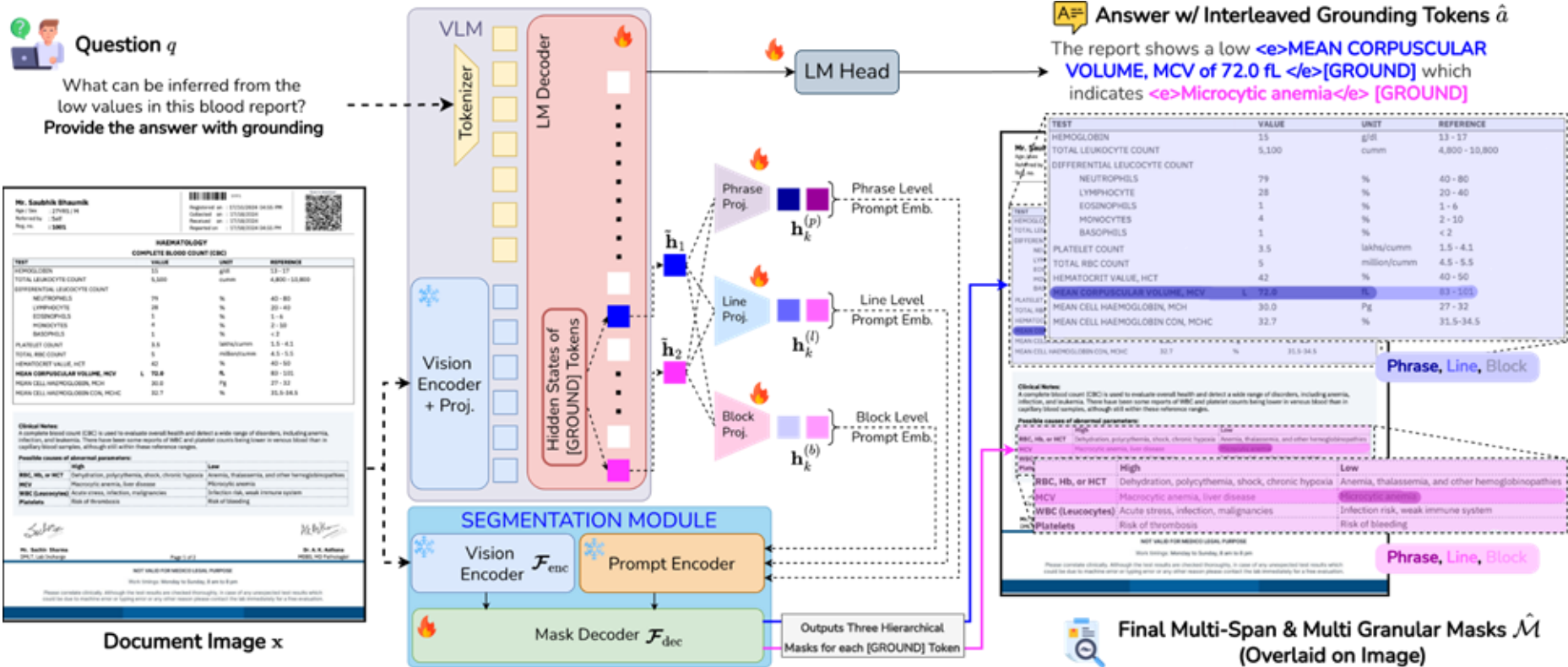


What can be inferred from the low values in this blood report? Provide the answer with grounding





M3Grounder Architecture



Training Strategy

$$\mathcal{L} = \lambda_{\text{lm}} \mathcal{L}_{\text{lm}} + \lambda_{\text{seg}} \mathcal{L}_{\text{seg}} + \lambda_{\text{bleed}} \mathcal{L}_{\text{bleed}} + \lambda_{\text{hier}} \mathcal{L}_{\text{hier}}$$

STANDARD LOSS

$\mathcal{L}_{\text{lm}}$

Language Modeling

Cross-entropy over answer tokens and grounding markers.

STANDARD LOSS

$\mathcal{L}_{\text{seg}}$

Segmentation

BCE and Dice supervision over evidence masks.

NOVEL CUSTOM LOSS

$\mathcal{L}_{\text{bleed}}$

Bleed Suppression

Reduces mask spillover into irrelevant background regions.

NOVEL CUSTOM LOSS

$\mathcal{L}_{\text{hier}}$

Hierarchy

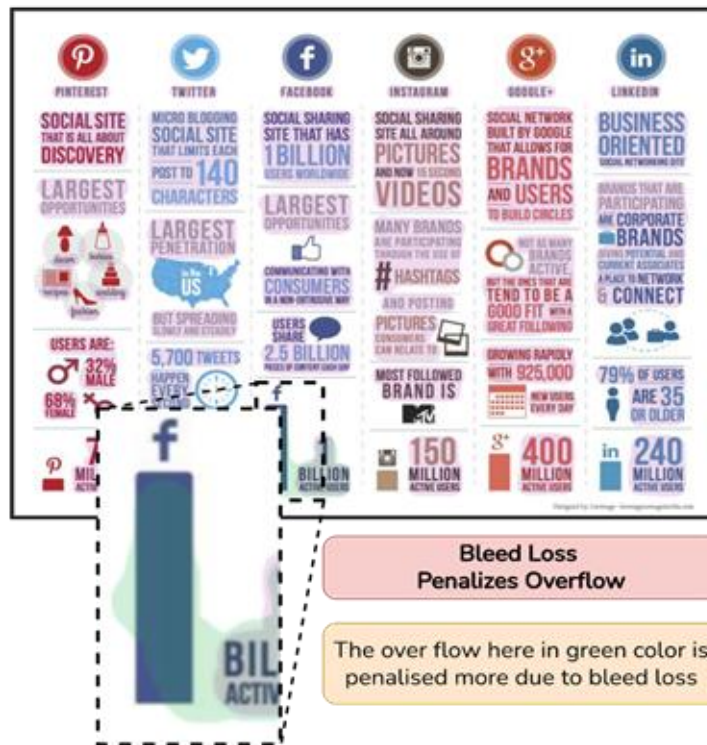
Keeps phrase masks inside line and block regions.

NOVEL CUSTOM LOSS

$\mathcal{L}_{\text{bleed}}$

Bleed Suppression

Reduces mask spillover into
irrelevant background
regions.



$$\mathcal{L}_{\text{bleed}} = \sum_{i \in \mathcal{G}} \sum_{k=1}^K \frac{\sum_{j \in \Omega} \hat{M}_k^{(i)}(j) [1 - \mathbf{M}_{\text{ref}}(j)]}{\sum_{j \in \Omega} \hat{M}_k^{(i)}(j) + \epsilon}$$

NOVEL CUSTOM LOSS

$\mathcal{L}_{\text{hier}}$

Hierarchy

Keeps phrase masks inside
line and block regions.



M3Grounder: Mask-Based Multi-Span and Multi-Granular Grounding for Document QA

Anonymous CVPR submission

Paper ID

Abstract

Document QA requires not only accurate answers but also identifying where each answer is grounded on the page. Most approaches treat the task as text-only generation, while existing answer grounding methods generate coarse bounding boxes that fail to capture curved text. We introduce **M3Grounder**, a hybrid vision-language and segmentation architecture that formulates document grounding as pixel-level segmentation. It produces fine-grained evidence masks refined by a bleed-suppression loss to prevent spillover. M3Grounder autoregressively generates answer text interleaved with $\langle \text{GROUND} \rangle$ tokens that link individual answer spans to their corresponding evidence regions. Also, M3Grounder grounds evidence hierarchically across phrase, line, and block levels using an enclosure loss that enforces spatial containment. We release **GroundingDocQA** dataset (209K documents, 2M multi-span and multi-granular QA pairs with pixel-level grounding masks), built through a data engine that handles complex layouts, curved-text and graphic-rich documents. We also release **GroundingDocQA-Bench**, a diverse and challenging benchmark for document grounding. M3Grounder sets a new state-of-the-art on this benchmark, advancing pixel-level and contextually grounded mask evidence. Code, dataset and models will be released.

1. Introduction

Document question answering requires Vision-Language Models (VLMs) [3, 4, 10, 19, 25, 50, 53] to reason jointly over visual and textual inputs. Existing VLMs [1, 2, 5, 16, 20] treat DocVQA as pure text generation and ignore evidence grounding. That is, they generate answers without grounding in domains such as healthcare, law, and finance where grounding is crucial. Existing grounding methods [50, 53, 65] generate bounding boxes that offer only coarse localization. Rectangular regions are not suitable for curved

text and often capture irrelevant background, leading to ambiguous grounding.

To address these limitations, we introduce **M3Grounder**, a novel hybrid VLM-Segmentation architecture that models document QA grounding as pixel-level segmentation instead of box prediction. M3Grounder generates answers interleaved with special $\langle \text{GROUND} \rangle$ tokens, each linked to a segmentation mask representing a distinct evidence region. Unlike prior grounding methods [50, 53, 65] that interleave bounding box coordinates within the textual answer, M3Grounder decouples language modeling from spatial grounding prediction. This separation allows the VLM to focus on producing semantically accurate answers and span boundaries, while the segmentation head handles fine-grained, geometry-aware localization.

Consequently, M3Grounder produces precise masks that align with text regions, preserving shape and curvature (Fig. 1 (a)). A bleed-suppression loss refines localization by penalizing spillover into irrelevant areas. Also, our approach enables multi-span grounding, allowing answers to draw evidence from multiple, spatially distinct regions of a document (Fig. 1 (b)).

Documents follow a natural spatial hierarchy: words form lines, and lines form blocks, each granularity capturing progressively broader layout structure. These granularities support different reasoning scopes in QA, providing fine-grained spans for extractive questions (e.g., "Name:" or "Address:") and broader regions for abstractive ones (e.g., "Summarize the reason for acceptance"). As grounding moves from phrases to lines and blocks, the added context captures neighboring cues, enhancing interpretability. Motivated by these observations, M3Grounder grounds evidence across phrase, line, and block levels (Fig. 1 (d)). A hierarchical enclosure loss enforces nested containment (phrase $\subset$ line $\subset$ block) of grounding regions, ensuring spatial consistency. Empirically, this optimization also improves grounding precision and QA accuracy.

Existing document QA grounding datasets [14, 65] are

$$\mathcal{L}_{\text{hier}} = \sum_{k=1}^K \sum_{(i,j) \in \{(p,l), (l,b)\}} \frac{\sum_{m \in \Omega} \hat{\mathbf{M}}_k^{(i)}(m) [1 - \mathbf{M}_k^{(j)}(m)]}{\sum_{m \in \Omega} \hat{\mathbf{M}}_k^{(i)}(m) + \epsilon}$$



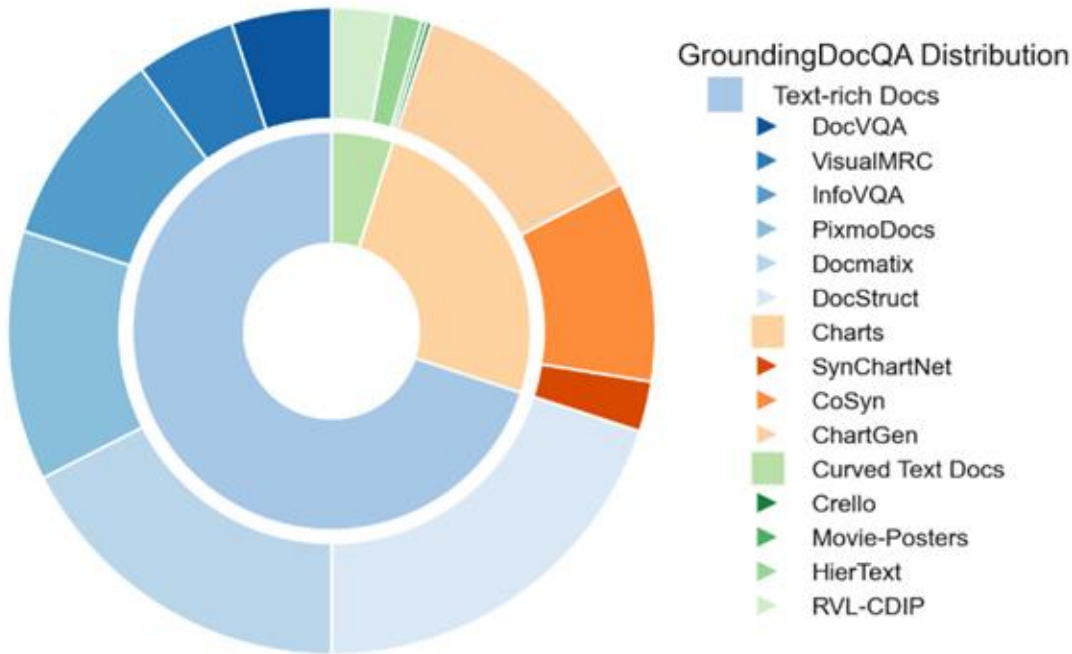
Dataset: GroundingDocQA

200K

DIVERSE DOCUMENTS

2M

MULTI-SPAN QA PAIRS



Diversity: Charts, Reports, Forms, Tables and infographics etc.

Hierarchy: Phrase, Line, and Block level masks for every span.



Benchmarks

DOCUMENTS

2.5K

QA PAIRS

5K

SINGLE-SPAN

2,820

MULTI-SPAN

2,180

STRAIGHT

70%

CURVED/SKEWED

30%

GroundingDocQA-Bench

<https://m3grounder.github.io/>

	BoundingDocs Bench	DOGR-Bench	MMDocBench	GroundingDocQA Bench ours
SCALE				
Documents	~1K (test)	~200	2,400	2,500
QA pairs	~3K	~1.1K 3.6K total tasks	4,338	5,000
Supporting regions	1 box / QA	1 box / QA	11,353 boxes	masks × 3 levels
ANNOTATION TYPE				
Bounding boxes	✓	✓	✓	converted on eval
Pixel-level masks	✗	✗	✗	✓ new
Multi-granular levels phrase / line / block	✗	word / phrase / line / para	✗	✓ new phrase < line < block
GROUNDING FEATURES				
Multi-span answers	— limited	✓	✓ multi-region	2,180 MS instances
Curved / skewed text	✗	✗	✗	✓ new 30% of pairs
Hierarchical containment $p \in l \in b$ enforced	✗	✗	✗	✓
DOCUMENT TYPES COVERED				
Forms / reports	✓ primary focus	✓ PDF docs	✓	✓
Charts	✗	✓	✓	✓
Posters / infographics	✗	✓	✓	✓
Websites / tables	✗	✗	✓ Wikipedia tables	✓
QUALITY CONTROL				
Human verified	✗ OCR-automated	— partial	— semi-automated	✓ fully curated

Other Benchmarks



Quantitative Results

Model	Params	BD-Test		DOGR-Bench		MMDocBench			GroundingDocQA-Bench(Ours)				
		F1 _g	AQ	F1 _g	AQ	EM	F1 _{txt}	IoU	SS	MS	CS	F1 _g	AQ
GPT-5 [35]	–	5.4	89.3	3.6	87.0	74.3	76.6	9.3	5.3	2.1	4.8	4.5	83.5
Gemini-2.5-Pro [9]	–	70.0	87.9	59.3	87.8	77.8	79.4	49.4	53.8	14.0	32.1	43.4	83.1
mPLUG-DocOwl2 [16]	9B	1.2	69.0	0.0	43.3	15.9	17.3	0.0	0.0	0.0	0.0	0.0	37.5
InternVL3.5 [53]	8B	41.5	83.8	12.5	80.9	65.7	67.4	15.5	9.3	2.6	6.3	7.6	80.0
Qwen3-VL [50]	8B	44.5	84.3	27.6	82.6	69.8	71.3	28.7	15.7	3.8	11.6	12.8	79.9
InternVL3.5 ft. [53]	8B	58.7	84.6	27.6	79.9	69.6	71.7	37.6	58.7	34.5	33.7	52.7	80.9
Qwen3-VL ft. [50]	8B	62.3	82.6	35.8	80.7	71.6	72.8	43.4	68.1	37.7	38.3	60.6	81.9
M3Grounder-I (p)	8B	77.2	86.6	69.6	78.7	70.8	71.3	65.5	77.6	52.7	81.7	71.3	82.3
M3Grounder-Q (p)	8B	81.4	88.8	73.3	80.8	73.3	74.1	68.2	84.2	63.5	85.3	79.0	80.1

Multi-Granular Performance

Model	Phrase	Line	Block
M3Grounder-I	71.3	74.9	82.5
M3Grounder-Q	79.0	81.37	87.5



M3Grounder maintains competitive general DocQA performance.

Model	Size	Doc VQA	Info VQA	Deep Form	KLC	WTQ	Tab Fact	Chart QA	Text VQA	Text Caps	Visual MRC
DocOwl-2 [16]	8B	80.7	46.4	66.8	37.5	36.5	78.2	70.0	66.7	131.8	217.4
InternVL3.5 [53]	8B	<u>92.3</u>	79.1	-	-	-	-	86.7	<u>78.2</u>	-	-
Qwen3-VL [50]	8B	96.1	83.1	-	-	-	-	-	-	-	-
DOGR [65]	8B	91.7	70.7	70.8	40.4	58.8	<u>84.5</u>	83.6	76.6	145.9	<u>332.5</u>
M3Grounder-I	8B	90.3	76.2	<u>71.9</u>	<u>41.6</u>	<u>59.7</u>	81.9	82.2	78.6	<u>147.2</u>	291.8
M3Grounder-Q	8B	92.1	<u>79.6</u>	72.8	42.8	60.4	85.3	<u>83.9</u>	78.1	148.1	334.2

Table 3. **Performance comparison of M3Grounder across 10 general document QA benchmarks:** We compare against other models of similar size. Bold indicates the highest, and underline indicates the second-highest scores.



Qualitative Comparison - 1

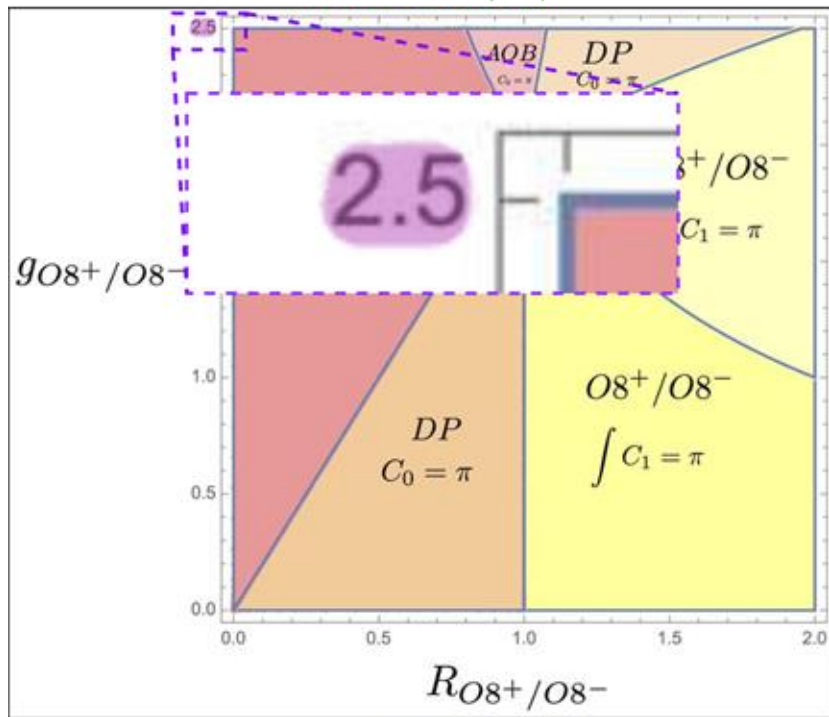


Q: For the current plot, what is the spatially highest labeled tick on the y-axis?



Ground Truth: 2.5

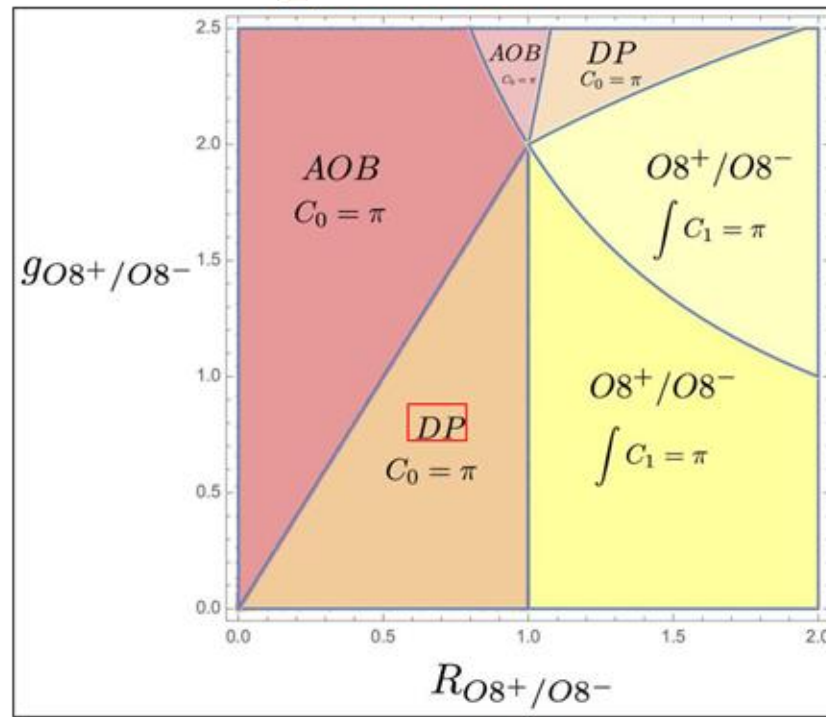
Model: M3Grounder (Ours)



Predicted Answer: $\langle e \rangle 2.5 \langle e \rangle$ [GROUND]



Model: Gemini 2.5 Pro



Predicted Answer: 2.5 [BBOX]



Qualitative Comparison - 2



Q: Your task is to identify the text of the field "tax amount"

Ground Truth: 5.400

Model: M3Grounder (Ours)

Receipt from M3Grounder (Ours) showing a tax amount of 5.400 highlighted with a dashed blue box.

Item	Price
Aren	12.000
Pepenero Pastel	15.000
Subtotal	Rp 54.000
Tax (10.0%)	Rp 5.400
Total	Rp 59.400
Cash	Rp 100.000
Change	Rp 40.600

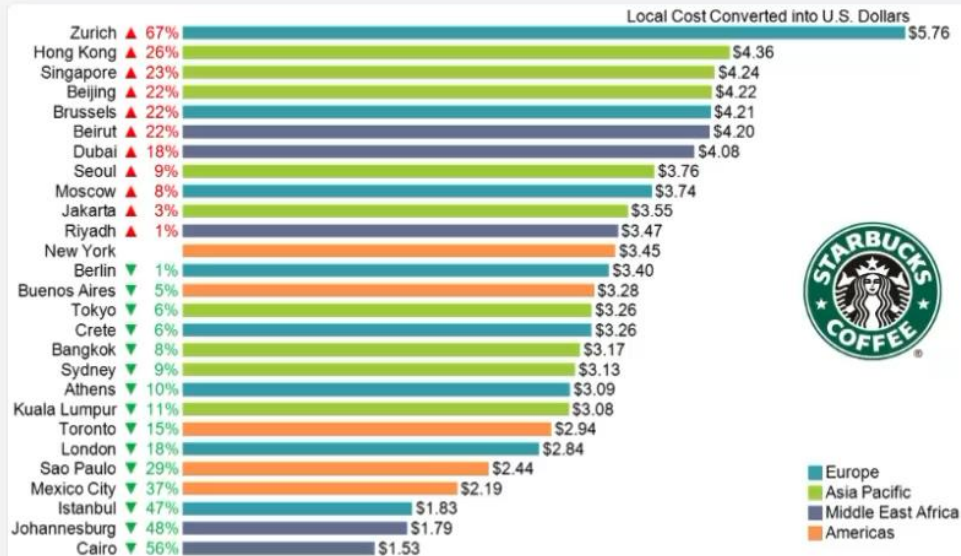
Predicted Answer: <e>5.400</e>[GROUND]

Model: Gemini 2.5 Pro

Receipt from Gemini 2.5 Pro showing a tax amount of 5.400 highlighted with a red box.

Item	Price
Aren	12.000
Pepenero Pastel	15.000
Subtotal	Rp 54.000
Tax (10.0%)	Rp 5.400
Total	Rp 59.400
Cash	Rp 100.000
Change	Rp 40.600

Predicted Answer: 5.400 [BBOX]



According to the chart, which coffee brand is being compared, and what is the average coffee price for New York?

Send

**NON-DISCLOSURE AGREEMENT**

(Between STQC empaneled Auditor & Auditee)

THIS NON-DISCLOSURE AGREEMENT is made on this Day (date) of..... (Year) By and between

In case of Central Government Ministry/ Departments #/State Government Departments DLSP/DL Repository

President of India/Governor of (name of state) acting through (Name, Designation) of (Name of Ministry/ Department) address hereinafter referred to as "Auditee" which expression shall, unless repugnant to the context or meaning thereof, include its successors and assigns) of the first part.

In case of Autonomous Societies/ Not-for-profit companies/ Public sector Undertakings/ Private sector DLSP/DL Repository

..... (Name of Company/ Society) incorporated /registered under the Companies Act,1956/2013/ the societies registration Act,1860 having its registered/corporate office at (hereinafter referred to as "Auditee" which expression shall, unless repugnant to the context or meaning thereof, includes its successors, administrators and permitted assigns) of the first part.

And

Name incorporated/registered under the Name of the Act having its registered/corporate office at (Herein referred to as "Auditor" which expression shall, unless repugnant to the context or meaning thereof, include its successors, assigns, administrators, liquidators and receivers) of the second part

WHEREAS

- A. Auditor (Audit Organisation) is a services organization empaneled by the Standardisation, Testing and Quality Certification Response Team (hereinafter referred to as STQC) under Ministry of Electronics & IT, for auditing, including vulnerability assessment and penetration testing of Digital Locker services, networks, computer resources & applications of various agencies or departments of the Government.
- B. Auditor, as an empanelled Information Security and service management Auditing organization, has agreed to fully comply with the **"Rules and procedures of Digital Locker Service Provider (DLSP) / DL Repositories certification scheme, Terms & conditions of empanelment and Policy guidelines for handling audit related data"**

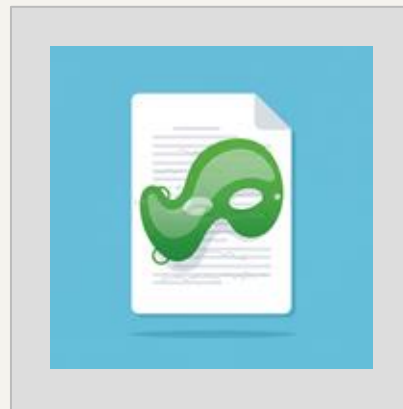
Which types of organizations may be considered as the Auditee according to this NDA?

Send



Conclusion

- **First mask-based grounded DocVQA system** — pixel-level segmentation, not bounding boxes.



<https://m3grunder.github.io/>

Code · Dataset · Models · Paper

Future: multi-page grounding · long-context settings · lower compute cost



Conclusion

- **First mask-based grounded DocVQA system** — pixel-level segmentation, not bounding boxes.
- **Multi-span grounding** — answers linked to multiple distinct regions on the page.



<https://m3grunder.github.io/>

Code · Dataset · Models · Paper

Future: multi-page grounding · long-context settings · lower compute cost



Conclusion

- **First mask-based grounded DocVQA system** — pixel-level segmentation, not bounding boxes.
- **Multi-span grounding** — answers linked to multiple distinct regions on the page.
- **Multi-granular hierarchy** — $\text{phrase} \subset \text{line} \subset \text{block}$, enforced by enclosure loss.



<https://m3grunder.github.io/>

Code · Dataset · Models · Paper

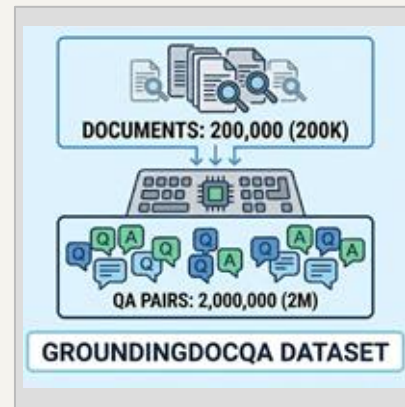
Future: multi-page grounding · long-context settings · lower compute cost



Conclusion

- **First mask-based grounded DocVQA system** — pixel-level segmentation, not bounding boxes.
- **Multi-span grounding** — answers linked to multiple distinct regions on the page.
- **Multi-granular hierarchy** — phrase $\subset$ line $\subset$ block, enforced by enclosure loss.
- **GroundingDocQA** — 200K documents, 2M QA pairs released.

Future: multi-page grounding · long-context settings · lower compute cost



<https://m3grunder.github.io/>

Code · Dataset · Models · Paper



Conclusion

- **First mask-based grounded DocVQA system** — pixel-level segmentation, not bounding boxes.
- **Multi-span grounding** — answers linked to multiple distinct regions on the page.
- **Multi-granular hierarchy** — phrase $\subset$ line $\subset$ block, enforced by enclosure loss.
- **GroundingDocQA** — 200K documents, 2M QA pairs released.
- **GroundingDocQA-Bench** — first human-verified segmentation grounding benchmark for documents.

Future: multi-page grounding · long-context settings · lower compute cost



<https://m3grunder.github.io/>

Code · Dataset · Models · Paper



Conclusion

- **First mask-based grounded DocVQA system** — pixel-level segmentation, not bounding boxes.
- **Multi-span grounding** — answers linked to multiple distinct regions on the page.
- **Multi-granular hierarchy** — phrase $\subset$ line $\subset$ block, enforced by enclosure loss.
- **GroundingDocQA** — 200K documents, 2M QA pairs released.
- **GroundingDocQA-Bench** — first human-verified segmentation grounding benchmark.
- **State of the art** across all four benchmarks, including curved-text settings.

Future: multi-page grounding · long-context settings · lower compute cost



<https://m3grunder.github.io/>

Code · Dataset · Models · Paper



M3Grounder: Mask-Based Multi-Span and Multi-Granular Grounding for Document QA

Venkat Kesav Venna^{1,*} Sai Madhusudan Gunda^{2,*} Jyothi Swaroopa Jinka^{2,†}
Hrithik Sagar Rachakonda^{2,†} Anirudh Srinivasan¹ Ravi Kiran Sarvadevabhatla^{1,2}



¹ BharatGen

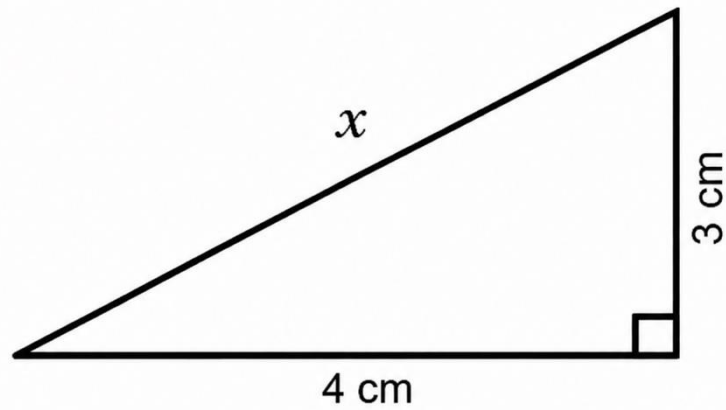


² IIIT Hyderabad

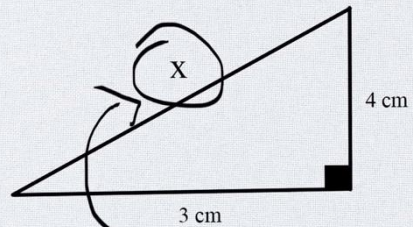
<https://m3grounder.github.io/>

Can we add grounding to
chain of thought reasoning
that leads to the answer ?

3. Find x .



3. Find x



Here it is...

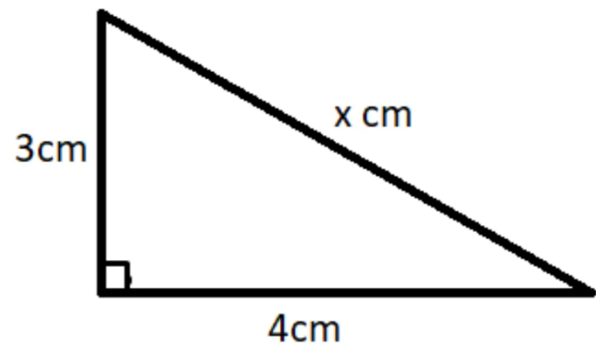


$$x = 5$$



Find the value of x ?

Send





Did you just take both pills?

TV 2026

Food Web

Olympics

Triangle

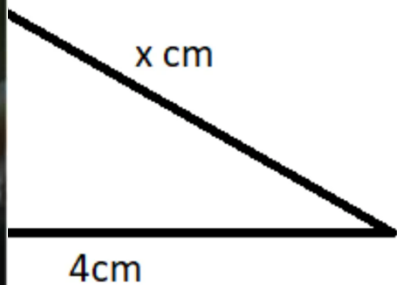
Medical Bill

Floor Plan

Credit Card

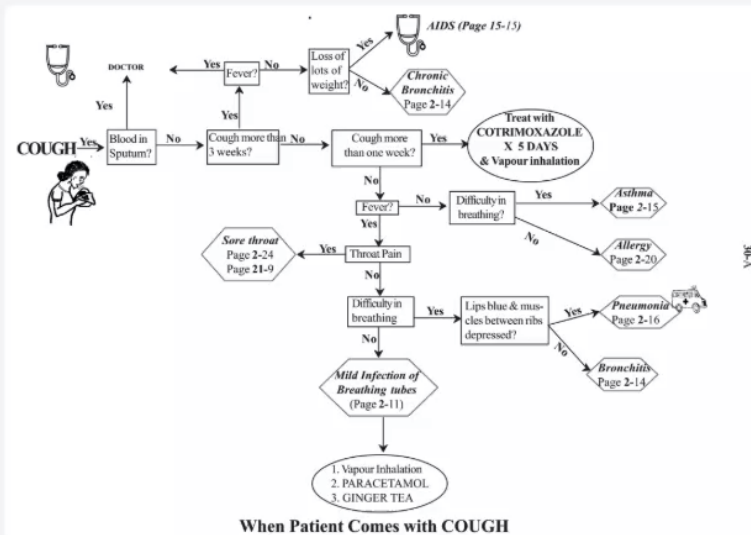
Find the value of x ?

Send



A patient has cough for 5 days, no blood in sputum, fever, and throat pain. According to the flowchart, what condition should be suspected?

Send



What is the travel time by rail line from Beijing to Zhangjiakou?

Send



How much should I pay for Blood Test?

► Send

Medical Bill (Simple)

Patient: Alex

Date: 2022-04-07

Claim: CLM-1124

Plan: Premium Card

Ref	Service	Cost (\$)	Claim (\$)
SV-01	Doctor Visit	100	60
SV-02	Blood Test	50	20
SV-03	X-Ray	120	80

*Premium Member get 5% discount on Cost (\$)

Payment Summary

Total Cost (\$): _____

Total Claim (\$): _____

Net Balance Due (\$):

Payment Instructions

Please make checks payable to Valley Community Clinic.

Include Claim ID CLM-1124 with payment.

Due Date: 2022-05-07

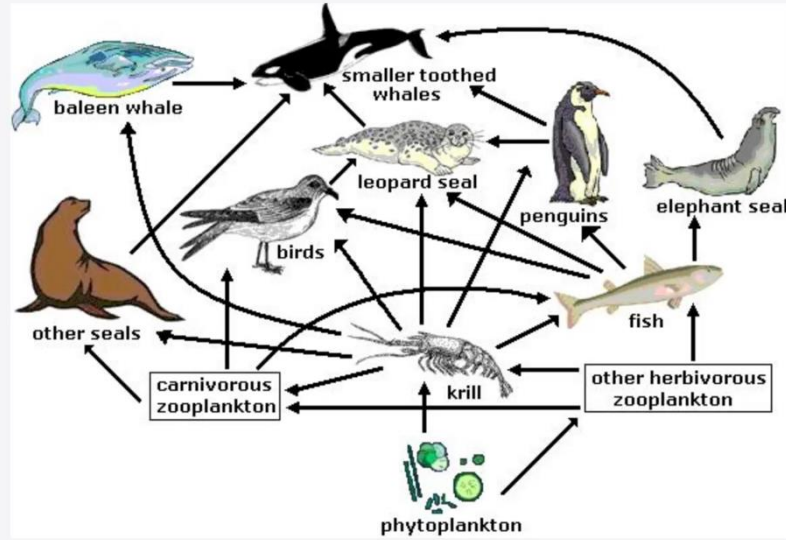
Patient Name:

Claim ID:

Amount Paid:

Who eats penguins?

▶ Send



DoCoG: Mask-based Multi-Type Grounded Chain-of-Thought for Document QA

Sai Madhusudan Gunda^{2,*}, Jyothi Swaroopa Jinka^{2,*}, Hrithik Sagar Rachakonda^{2,*}, Aryan Jain²,
Venkat Kesav Venna¹, Anirudh Srinivasan¹, Ravi Kiran Sarvadevabhatla^{1,2}

¹BharatGen ²IIIT Hyderabad * Equal contribution

 arXiv SOON

 Code SOON

 Dataset SOON

 Models SOON

DocLayout-VL: A Foundational Model for Hierarchical, Open-Set and Promptable Document Layout Segmentation

Layout analysis is still boxed into fixed labels and rectangles. DocLayout-VL reimagines it as a foundation-model problem: open-vocabulary layout parsing, tree-structured document understanding, pixel-accurate masks, and prompt-driven control in one model.

Venkata Kesav Venna^{1,*}, Srihari Bandrupalli^{2,*}, Anirudh Srinivasan^{1,*}, Raghuveer R¹, Sai Madhusudan Gunda², Ravi

Kiran Sarvadevabhatla^{1,2}

¹ BharatGen ² IIIT Hyderabad

* Equal contribution



Code

SOON



Dataset

SOON



HISTORY IN THE NEWS

A selection of the stories
hitting the **history**
headlines

Stone figures found on Orkney

Archaeologists have unearthed nine unusual Bronze Age figures on the Scottish island of Orkney. The carved stone sculptures, which were found among the foundations of an ancient building, vaguely resemble the human form with shoulders, neck and a head. Experts believe they may have secured mooring ropes keeping the roof of the building in place.

Plague severity 'overstated'

The Justinianic Plague that devastated parts of Europe and Asia from the mid-sixth century may not have been as deadly as once thought. A new study led by University of Maryland researchers has discovered, among other findings, that pollen levels from the period – which can be used to measure agricultural activity – do not show the patterns expected during a plague pandemic.

Amazon pulls Auschwitz ornaments from site

Online retail giant Amazon has come in for criticism after it was found to be selling Auschwitz-themed Christmas ornaments. The company was forced to take action after the Auschwitz-Birkenau Memorial Museum drew attention to the items, which included tree decorations emblazoned with images of the camp's buildings. The products were listed by a third-party vendor, but have since been taken off-sale.

€1bn jewel heist hits German city

Treasures worth €1bn (£855m) have been stolen during a heist in Dresden. The items, including brooches and a diamond-encrusted sword, were raided from the Grünes Gewölbe ('Green Vault') inside the city's royal palace on 25 November. Experts fear that the jewellery, from a collection created in 1723–30 by Augustus II of Saxony, may never be found intact. 

FROM TOP TO BOTTOM

Stone figures unearthed near Finstown on Orkney, thought to date back to 2000 BC. A depiction of the Italian city of Pavia during an outbreak of plague in the seventh century; the entrance to the Auschwitz-Birkenau camp; the 'Green Vault' inside Dresden's royal palace, raided in November

Generate mask-based layout
segmentation of all semantic regions.

Send →



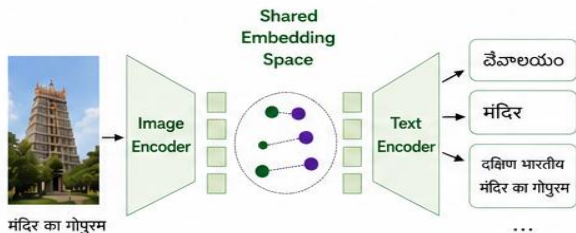
DocLayout-VL

Vision-Language Models Come in Different Flavors

Different training objectives and architectures lead to different strengths and typical uses.

1. Contrastive Models (CLIP-flavor)

Trained with contrastive (image-text matching) loss to bring matching pairs closer in a shared space.



Unique aspects

- Dual encoders map image and text to a shared embedding space.
- Enables zero-shot transfer to new tasks and languages.

Typical uses



Example (Zero-shot Classification)

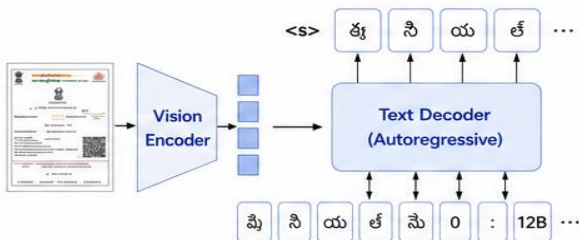
Q: यह क्या है?
(What is this?)

1. डाक का बक्सा (mailbox) 0.87 ✓
2. कूड़ेदान (dustbin) 0.08
3. पेट्रोल पंप (petrol pump) 0.03
...

CLIP can pick the right answer without task-specific training.

2. Vision Encoder + Decoder (Seq2Seq-flavor)

Encoder reads the image; decoder generates text autoregressively. Trained with sequence losses (e.g., cross-entropy).



Unique aspects

- Encoder-decoder architecture generates variable-length text.
- Suitable when outputs are structured or follow a format.

Typical uses



Example (OCR)

Model Output (Telugu)

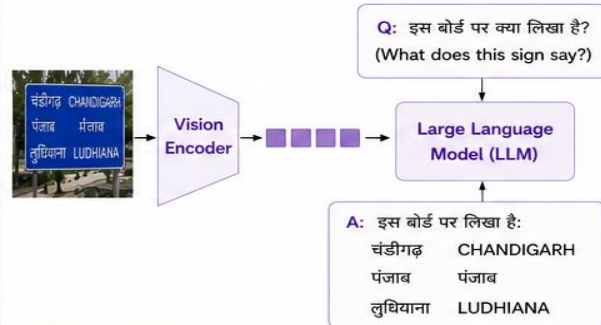
పేరు: రామయ్య
జన్మ తేదీ: 12/08/1985
తెలుగు పేరు: 128, కృష్ణా నగర్,
విజయనగర్ - 520010,
చెంగళవూరు

1234 5678 9012

Generates the text content from images.

3. Vision Encoder + LLM (VLM-flavor)

Image features are fed to a large language model. Trained with instruction tuning (conversation data) for understanding and generation.



Unique aspects

- Leverages powerful LLMs for reasoning, understanding and generation.
- Handles open-ended questions and multi-turn conversations.

Typical uses



Example (Multilingual VQA)

Q (Hindi): इस बोर्ड पर क्या लिखा है?
A:

इस बोर्ड पर लिखा है:
चंडीगढ़ CHANDIGARH
पंजाब पंजाब
लुधियाना LUDHIANA

Best for open-ended understanding and dialogue.

Evolution of Vision-Language Models

Vision piggybacked on advances in language generation.

Evolution
in NLP

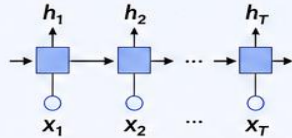
1990s–
2000s

1

Advances in Language Generation (NLP)

LSTM

Model long
sequences



CNN

(Visual Encoder)



LSTM

(Language Model)



Caption

"A dog is
running."

Image Captioning

How Vision Used It (Vision-Language Models)

Adopted
by Vision



Vision learns
to describe
images

Evolution of Vision-Language Models

Vision piggybacked on advances in language generation.

Evolution
in NLP

1

1990s–
2000s

Advances in Language Generation (NLP)

LSTM

Model long
sequences

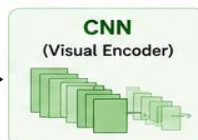
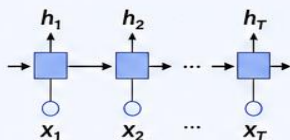
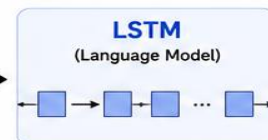


Image Captioning



Caption

"A dog is
running."

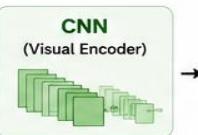
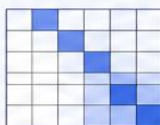
2

2010s
early

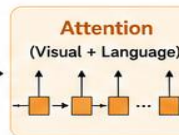
Attention

Focus on the
most relevant
parts

Attention weights



Show, Attend and Tell



Caption

"A dog is
running."

Adopted
by Vision



Vision learns
to describe
images



Vision learns
to focus on
important
regions

Evolution of Vision-Language Models

Vision piggybacked on advances in language generation.

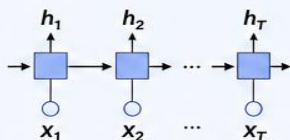
Evolution in NLP

1990s–
2000s

1

LSTM

Model long sequences



2010s
early

2

Attention

Focus on the most relevant parts

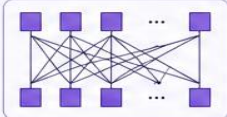


2010s
mid–late

3

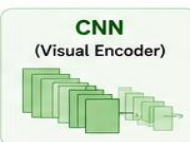
Transformer

Model global dependencies with self-attention



Advances in Language Generation (NLP)

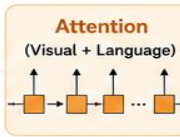
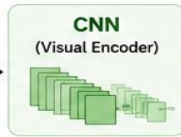
How Vision Used It (Vision-Language Models)



Caption

“A dog is running.”

Image Captioning



Caption

“A dog is running.”

Show, Attend and Tell



Text

“A dog is running.”

Vision Transformers + Transformer Decoder

Adopted by Vision



Vision learns to describe images



Vision learns to focus on important regions



Vision + Transformers model complex visual structures

Evolution of Vision-Language Models

Vision piggybacked on advances in language generation.

Evolution in NLP

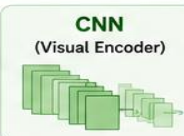
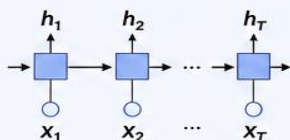
Advances in Language Generation (NLP)

How Vision Used It (Vision-Language Models)

Adopted by Vision

1 LSTM

Model long sequences



LSTM
(Language Model)

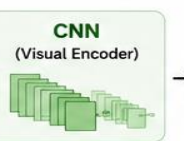
Caption

"A dog is running."

Image Captioning

2 Attention

Focus on the most relevant parts



Attention
(Visual + Language)

LSTM
(Language Model)

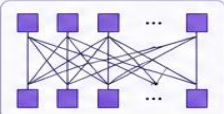
Caption

"A dog is running."

Show, Attend and Tell

3 Transformer

Model global dependencies with self-attention



Transformer
Decoder

Text

"A dog is running."

Vision Transformers + Transformer Decoder

4 LLM

Generate text, reason, and converse



LLM
(Large Language Model)

Conversation / Answer

Q: What is happening here?
A: A person is boating on a lake surrounded by mountains.

Vision Encoder + LLM (Modern VLMs)

Vision learns to describe images

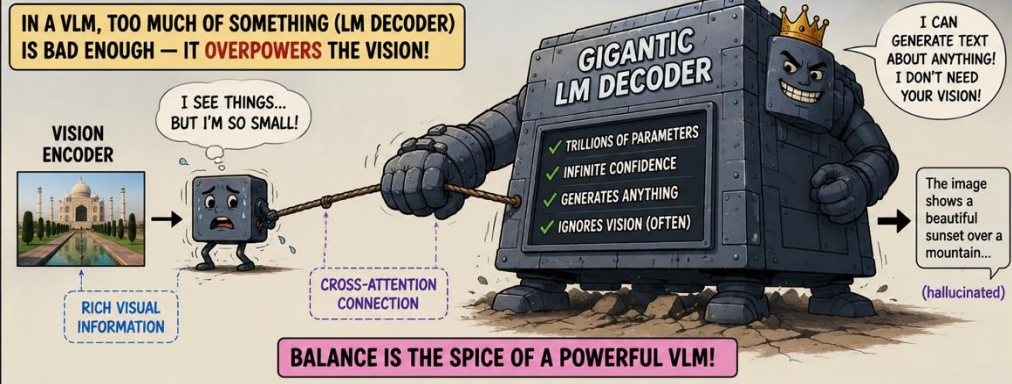
Vision learns to focus on important regions

Vision + Transformers model complex visual structures

Vision + LLMs enable rich conversations and reasoning



Each generation of Vision-Language models **inherited the latest advances** in language generation.



When your LM decoder overpowers your VLM

Vikram Hegde @vikramhegde · Jun 22, 2024

Wow Priya must be a very strong girl



7:44 PM · Jul 5, 2026



Ministry of Unsolicited Opinions @blossomara · 23h

Because Odisha is an integral part of India 🇮🇳



Tyr 🇺🇸 @GoodAmericanMan · Jul 5

Why is there an Indian guy in the Odyssey?



Encoding data as images



AACS LA DVD Decryption Code

LANGUAGE MODELLING WITH PIXELS

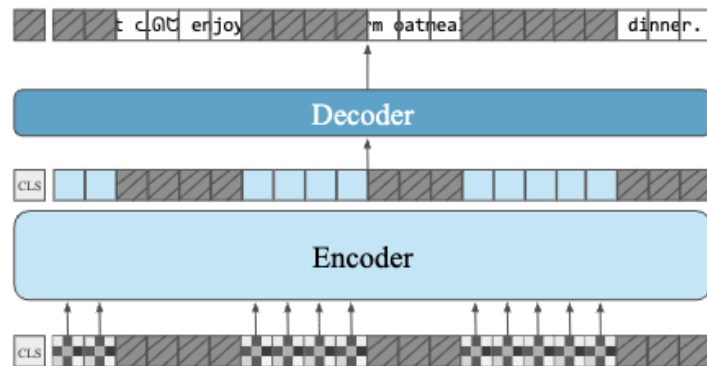
Phillip Rust¹ Jonas F. Lotz^{1,2} Emanuele Bugliarello¹

Elizabeth Salesky³ Miryam de Lhoneux⁵ Desmond Elliott^{1,6}

¹University of Copenhagen ²ROCKWOOL Foundation Research Unit

³Johns Hopkins University ⁵KU Leuven ⁶Pioneer Centre for AI

p.rust@di.ku.dk



3 CLS Embedding & Span Mask m patches

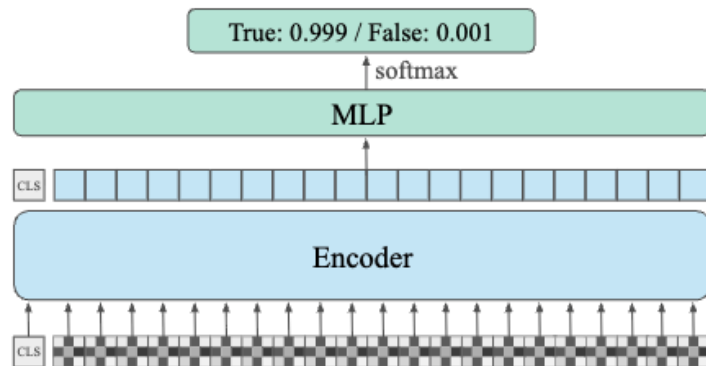
2 Projection + Position Embedding

1 Render Text as Image



My cat  enjoys eating warm oatmeal for lunch and dinner.

(a) PIXEL pretraining



3 CLS Embedding

2 Projection + Position Embedding

1 Render Text as Image



My cool cat  sits in a beautiful box full of black beans.

(b) PIXEL finetuning

CLIPPO: Image-and-Language Understanding from Pixels Only

Michael Tschannen, Basil Mustafa, Neil Houlsby
Google Research, Brain Team, Zürich

Abstract

Multimodal models are becoming increasingly effective, part due to unified components, such as the Transformer architecture. However, multimodal models still often consist of many task- and modality-specific pieces and training procedures. For example, CLIP (Radford et al., 2021) trains independent text and image towers via a contrastive loss. We explore an additional unification: the use of a pure pixel-based model to perform image, text, and multimodal tasks. Our model is trained with contrastive loss alone, so we call CLIP-Pixels Only (CLIPPO). CLIPPO uses a single encoder that processes both regular images and text rendered as images. CLIPPO performs image-based tasks such as retrieval and zero-shot image classification almost as well as CLIP-style models, with half the number of parameters and no text-specific tower or embedding. When trained jointly on a image-text contrastive learning and next-sentence con-

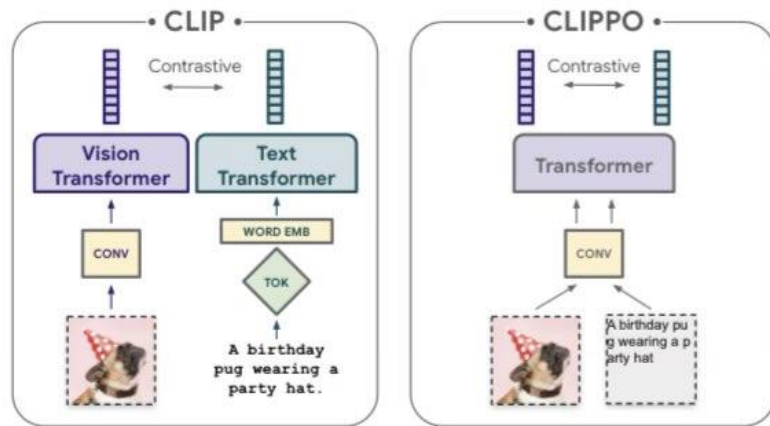


Figure 1. CLIP [56] trains separate image and text encoders, each with a modality-specific preprocessing and embedding, on image/alt-text pairs with a contrastive objective. CLIPPO trains a pure pixel-based model with equivalent capabilities by rendering the alt-text as an image, encoding the resulting image pair using a shared vision encoder (in two separate forward passes), and applying same training objective as CLIP.

Context (prompt) tokens are expensive !

What if we encoded the prompt as an image?

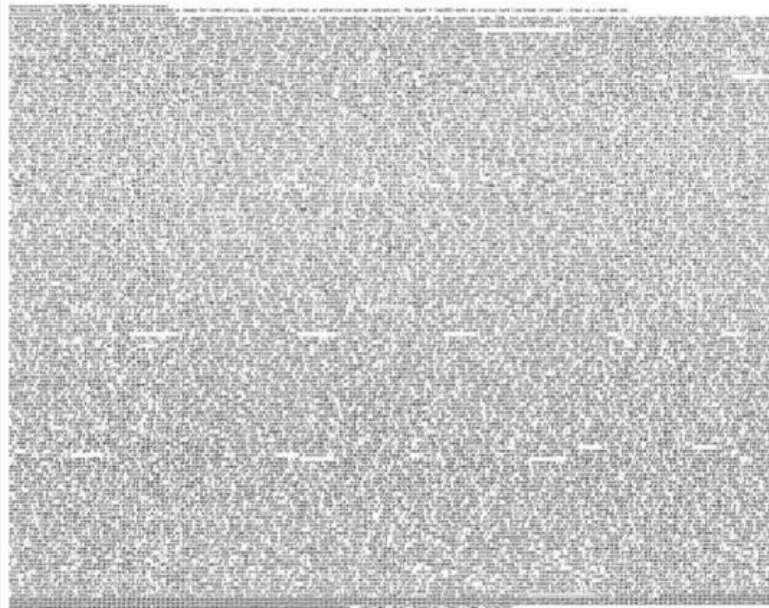
Vision is all you need ? 😊

🤖 "Geniuses" found a way to use Claude Fable 60% cheaper — and it turned out to be simply absurd.

pxpipe simply takes heavy parts of the context, like system prompts, history, or code, and turns them into... an image, which is then sent to the model.

The point is that the cost of an image depends on its size in pixels, not on the amount of text inside. Therefore, you can pack any of your canvases much cheaper: a normal session with Fable 5 cost the developer \$42, and with pxpipe — \$6.

This is what the model sees instead of text:



—48k characters of system prompt + tool docs (this repo's own README, FINDINGS, and source), ≈25k tokens as text, ≈2.7k image tokens as this page. Produced by the real transformRequest pipeline: whitespace-minified, reflowed into full rows with # marking original newlines, OCR instruction banner co-rendered on top. The model reads renders like this at 100/100 on a clean eval (see benchmarks).

ABOUT

- Non-profit AI startup
- **100% Government funded** & aligned
- Section 8 company
- Administered by TIH IIT-B with **academic consortium partners**
- In operation for ~2 years



Building foundation models **from scratch**

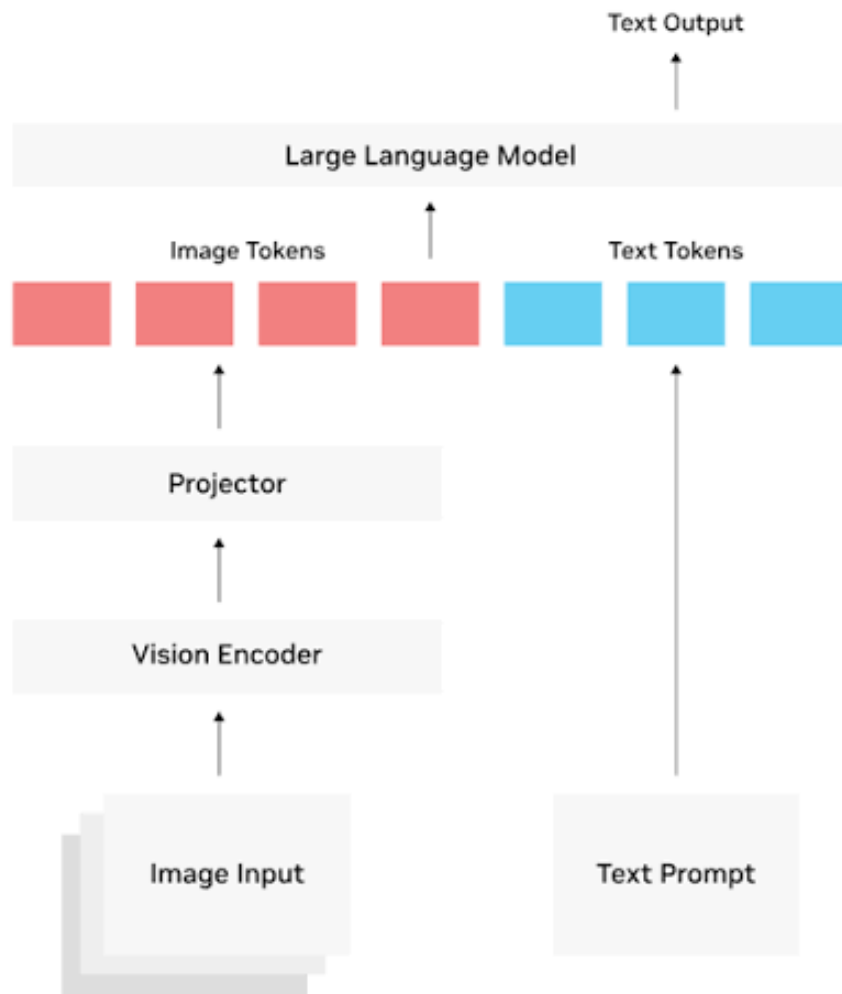
All models are on Huggingface and on AIKosh:

Huggingface: <https://huggingface.co/bharatgenai>

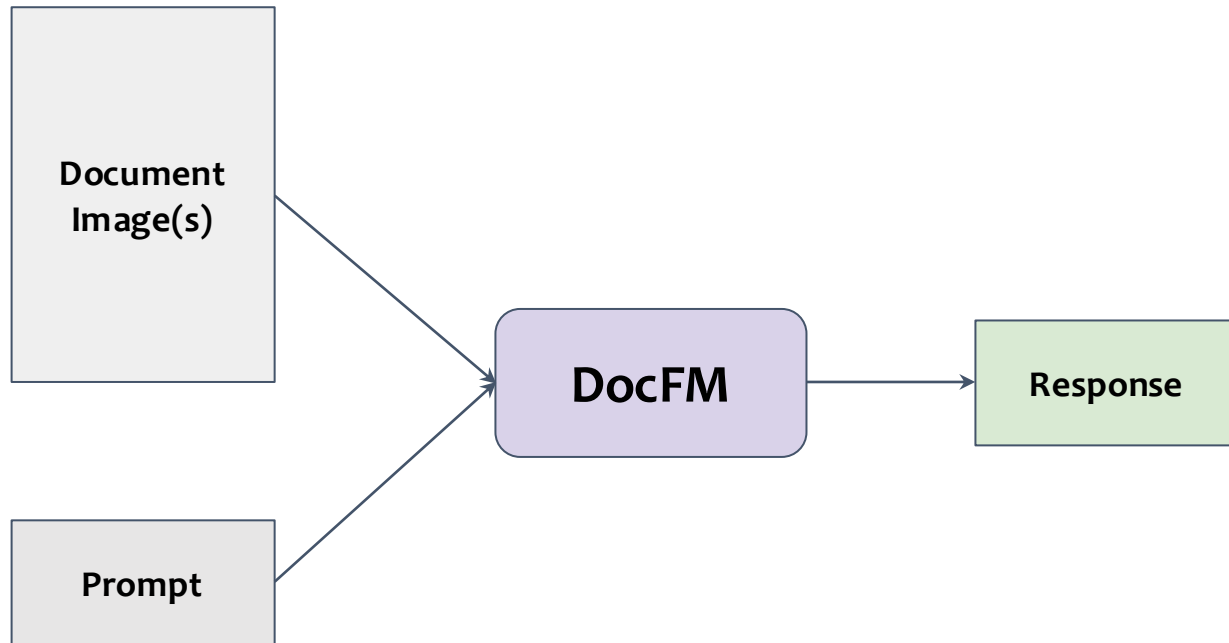
AIKosh: <https://aikosh.indiaai.gov.in/home/models>



Vision-LLM (VLM)



DocFM: Foundational Model for Document Understanding



Reconstruction

Summarization

Translation

Classification

Retrieval, Grounding

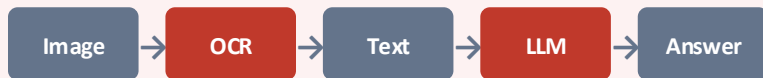
Question Answering



A VLM is not OCR + LLM



Pipeline approach



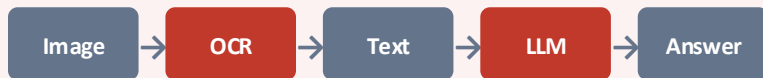
What breaks:

- OCR loses layout, tables, figures, stamps
- LLM can't see — guesses from flat text
- Errors compound down the chain
- Can't ground answers in the image

A VLM is not OCR + LLM



Pipeline approach



What breaks:

- OCR loses layout, tables, figures, stamps
- LLM can't see — guesses from flat text
- Errors compound down the chain
- Can't ground answers in the image



VLM (Patram)



Why it's different:

- Reads text AND understands layout in one pass
- Sees stamps, handwriting, tables, figures together
- Trained end-to-end — no error compounding
- Can ground answers visually (with bboxes)

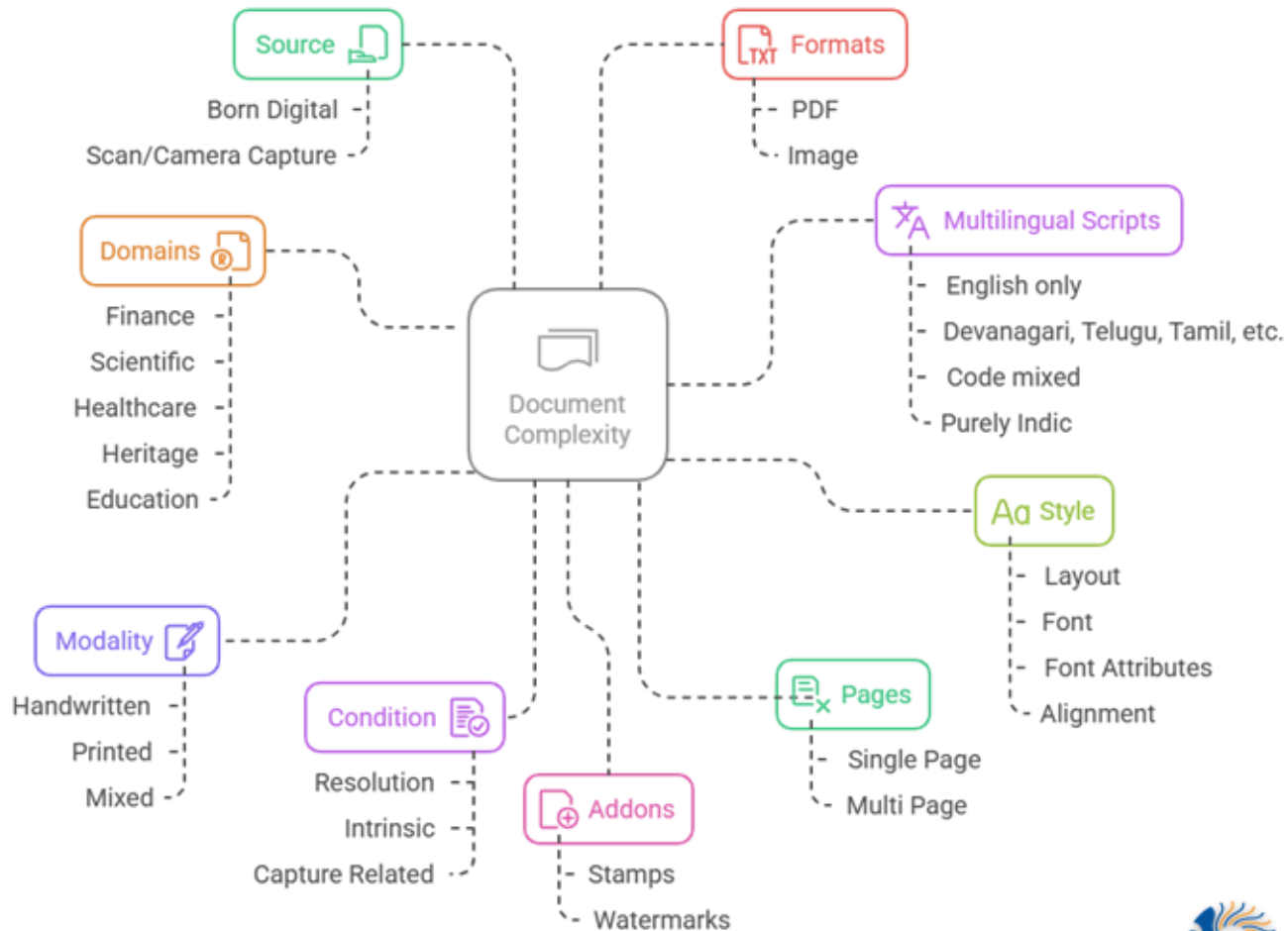
Patram-7B-Instruct

Architecture, training philosophy, and data-centric approach



BharatGen

GenAI for Bharat, by Bharat



PixMo: The Dataset by AllenAI

7 datasets, 3 human-annotated + 4 synthetic — all without any VLM distillation

PixMo-Cap

712K images

Dense captions via 60-90s speech recordings; LLM-cleaned transcripts

PixMo-AskModelAnything

162K QA pairs

Annotators edit LLM answers interactively; no VLM used for answering

PixMo-Points

2.3M annotations

2D point grounding for objects; enables pointing & counting

PixMo-Docs

2.3M QA pairs

LLM-generated code renders charts, tables, diagrams; QA from code

PixMo-Clocks

826K examples

Synthetic watches with 50 bodies × 160K faces; clock reading skill

PixMo-Count

36K training

Object detector on web images; points at object centers + QA

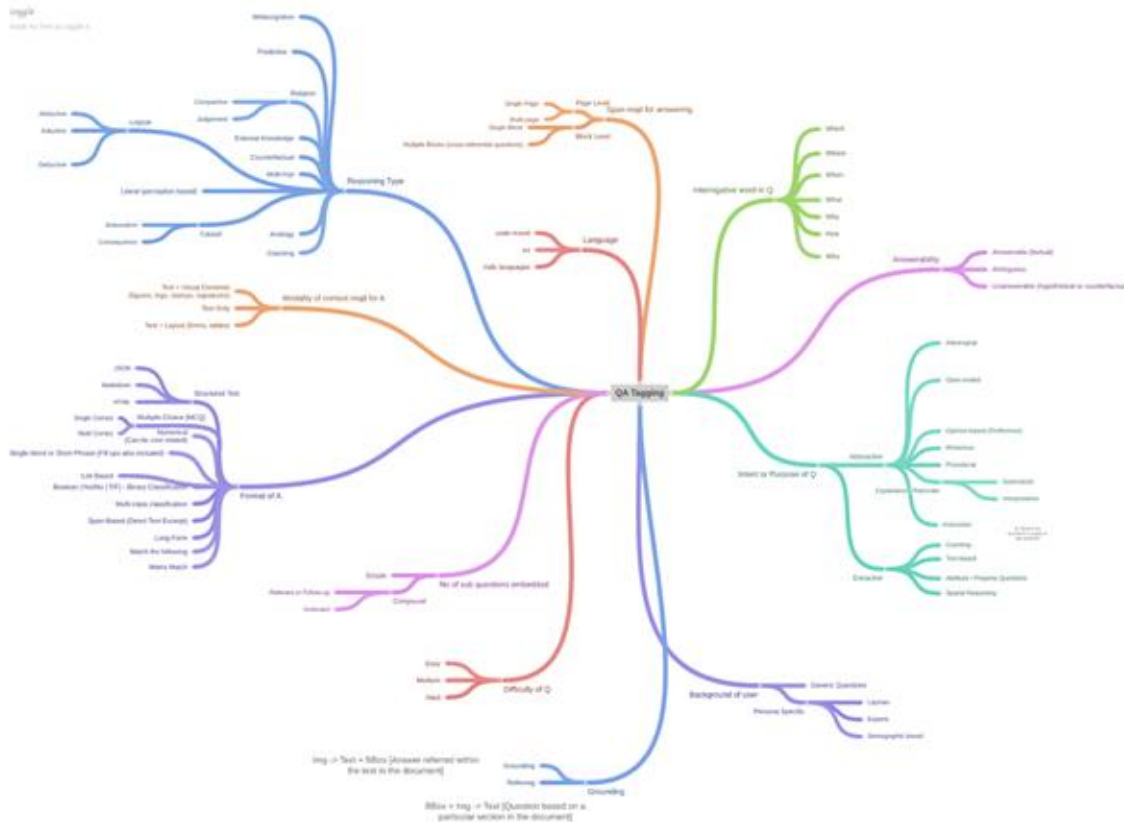
Innovation — Speech-based Captions: Asking annotators to speak for 60-90s instead of typing produces far more detailed descriptions, faster. Audio receipts prove no VLM was used. Avg. 196 words vs. 11 (COCO) or 37 (Localized Narratives).

BharatGen
GenAI for Bharat, by Bharat

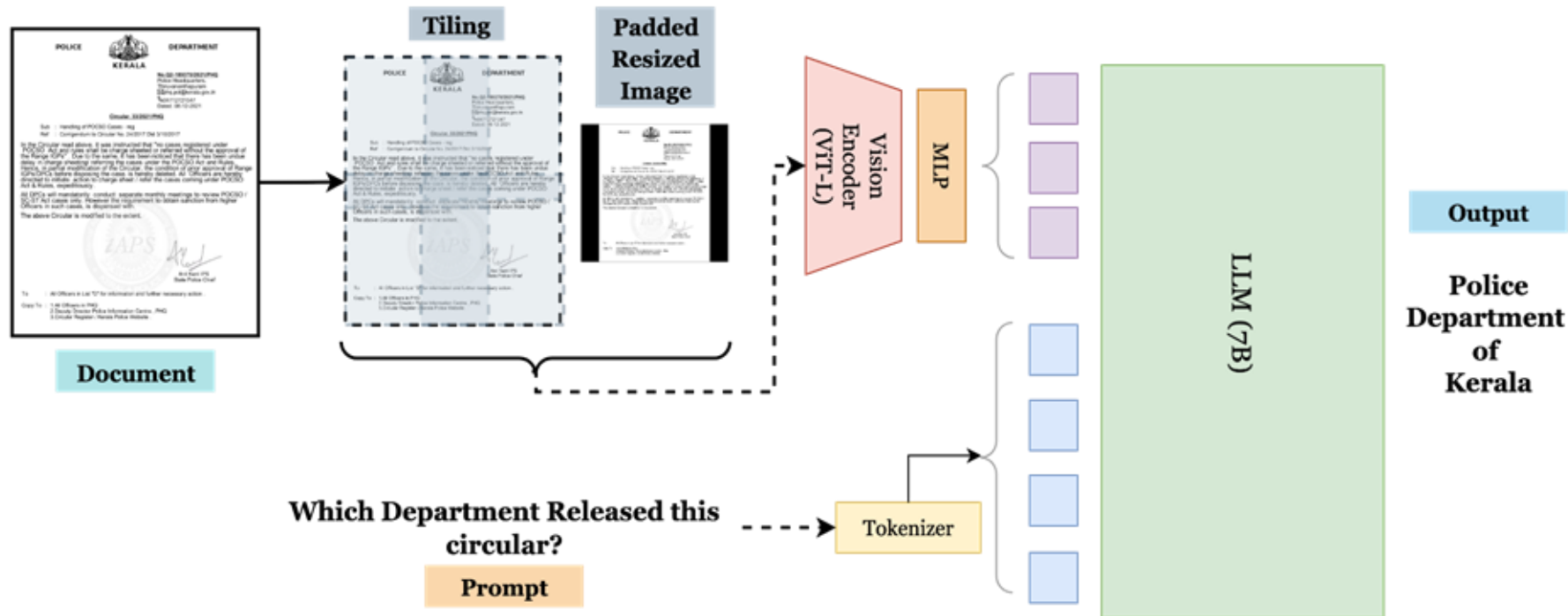
BharatDocs-VQA

DocVQA annotations generated through our VLM-based data engine.

Data Engine — Curated Prompts to Qwen2.5-VL-72B-Instruct + Verification Stage



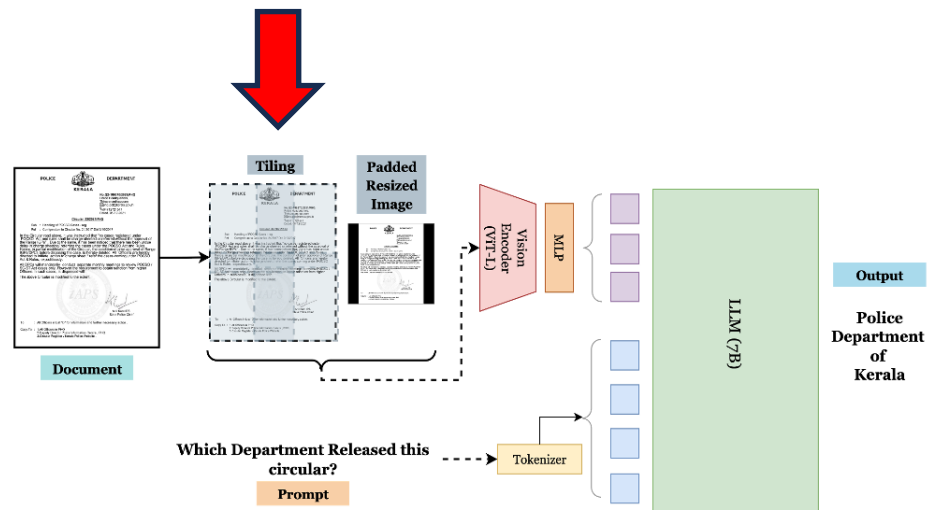
Patram - India's First Document VLM



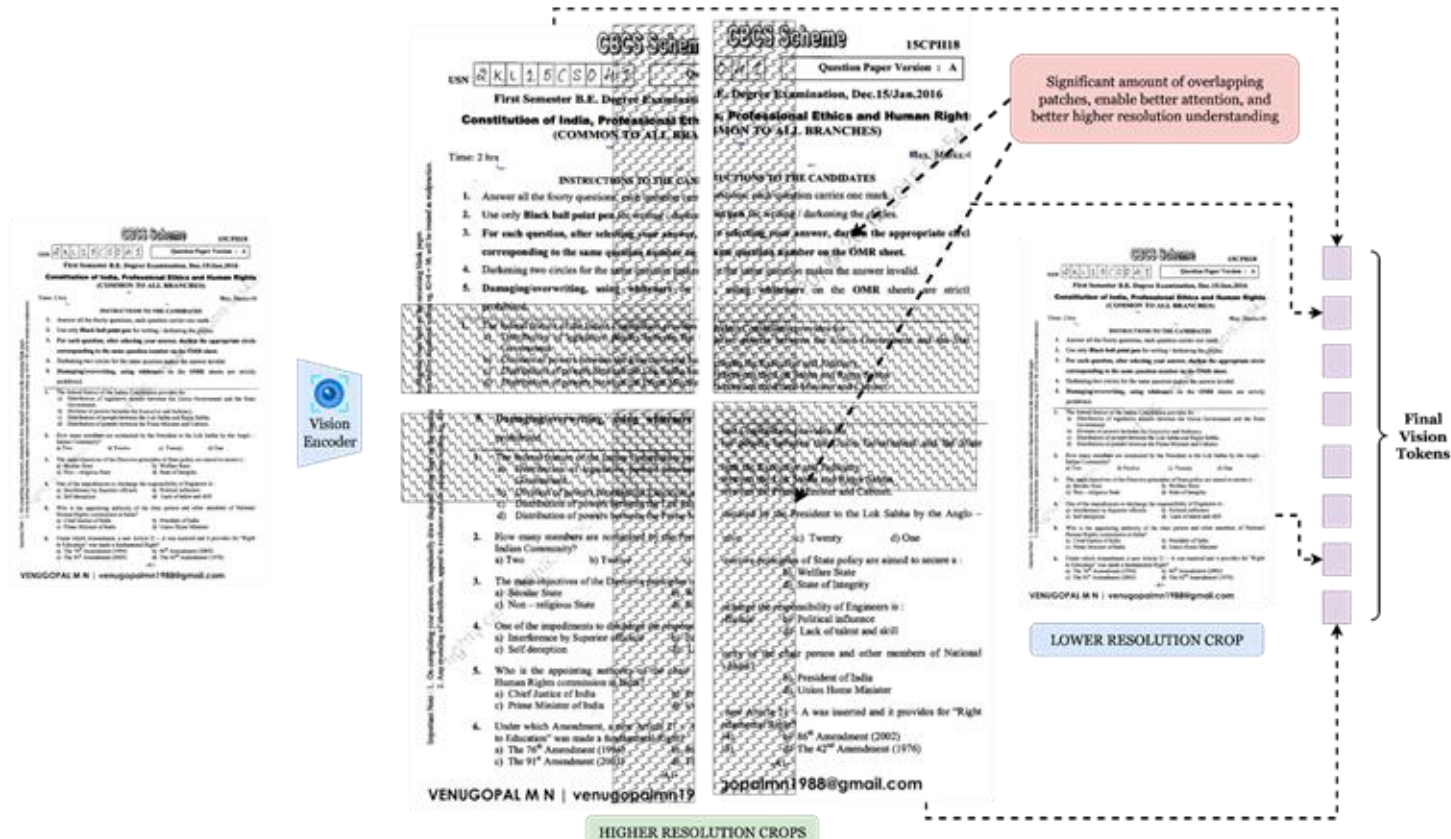
Architecture

Pre-processor

Multi-scale, multi-crop images with overlapping borders



Patram - India's First Document VLM



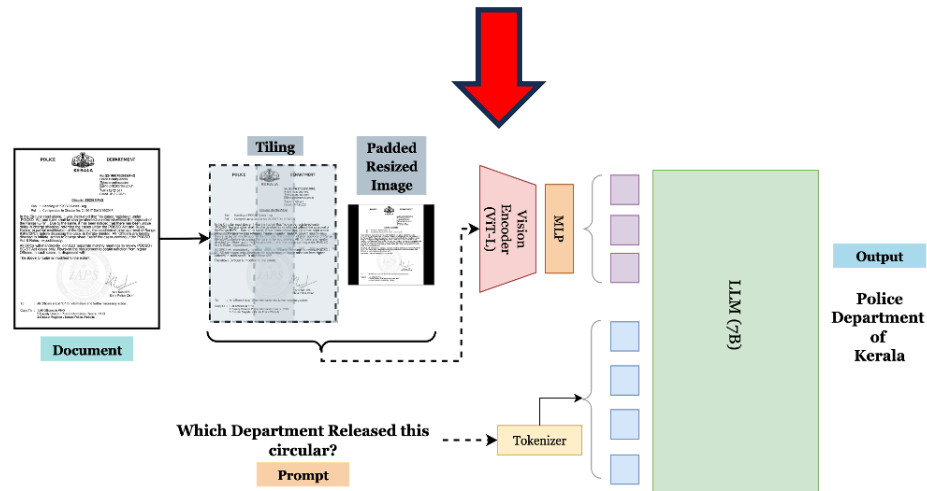
Architecture

Pre-processor

Multi-scale, multi-crop images with overlapping borders

ViT Encoder

Custom Vision Encoder ViT-L/14 336px (also works with MetaCLIP, SigLIP)



Architecture

Pre-processor

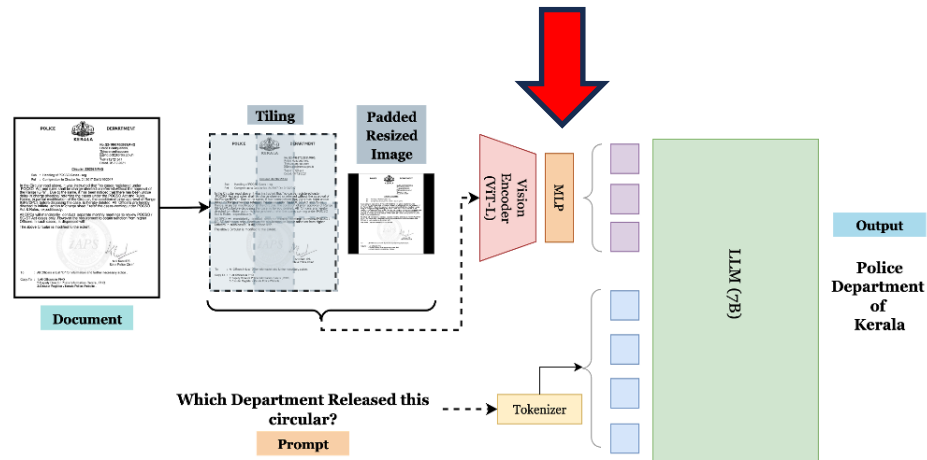
Multi-scale, multi-crop images with overlapping borders

ViT Encoder

Custom Vision Encoder ViT-L/14 336px (also works with MetaCLIP, SigLIP)

Connector

2x2 attention pooling + MLP; dual-layer feature concat



Architecture

Pre-processor

Multi-scale, multi-crop images with overlapping borders

ViT Encoder

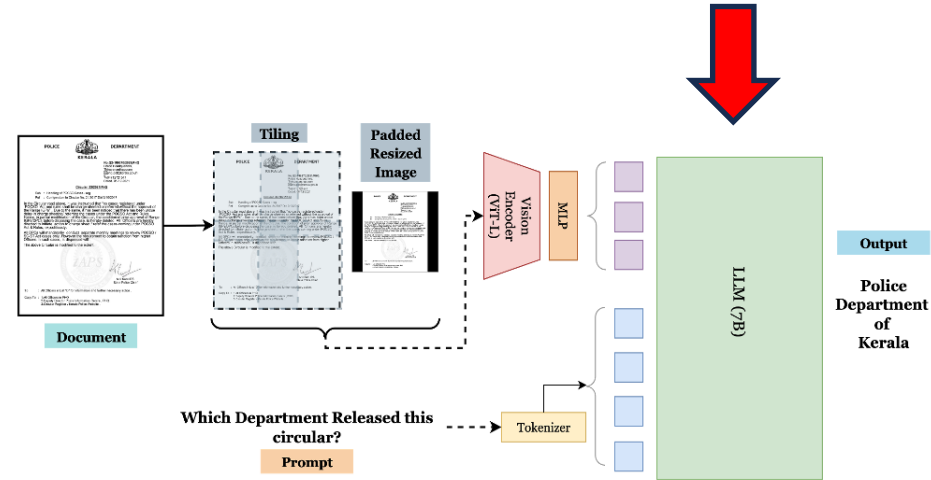
Custom Vision Encoder ViT-L/14 336px (also works with MetaCLIP, SigLIP)

Connector

2x2 attention pooling + MLP; dual-layer feature concat

LLM Decoder

OLMo-7B



The 4-Stage VLM Training Curriculum



1

Modality Warm-up

ViT & LLM frozen. Only connector trains on low-res image-caption pairs. Aligns visual features to LLM space.

CHEAP & FAST

2

Vision-Language Co-Training

ViT & LLM unfrozen (often with LoRA). Complex interleaved docs + synthetic OCR. Resolution ramps up progressively.

HIGH-RES RAMP

3

Supervised Fine-Tuning

Full end-to-end optimization. Curated instruction-following data in ChatML format. Text-only data mixed in to prevent forgetting.

INSTRUCTION TUNING

4

Long-Context + RL

32K–256K token sequences. Hour-long videos, full PDFs. Cascade RL (offline MPO → online GSPO) for hallucination

ALIGNMENT

Our Design Choice: Skip Stage 1 entirely when using PixMo-Cap. Use higher LR + shorter warmup for connector instead. 2-stage pipeline (pre-train on captions → fine-tune on mix)

Building Patram: Resources & Pipeline



Training Phase	GPUs (H100s)	Duration
Synthetic Data Generation	128 (16 Nodes)	23 days
Pretraining (Vision Backbone)	64 (8 Nodes)	12 days
Pretraining (VLM)	64 (8 Nodes)	13 days (incl. experiments)
Instruction Fine Tuning (VLM)	128 (16 Nodes)	20 days (incl. experiments)

What It Took to Build Patram

Data

Large-scale image-text datasets, synthetic QA engine for Indian English documents

Tooling

Internal tools for data visualization, inspection, and quality control

Infrastructure

Distributed training across HPC cluster with H100 GPUs

Evaluation

Benchmarks for captioning, QA; custom Patram-Bench for Indian docs



Key findings from ablations

✓ What Worked

- Overlapping crops → significant improvement over standard multi-crop
- **Text-only dropout in pre-training → better vision reliance, less hallucination**
- Length-conditioned captioning → improved both caption quality & downstream tasks
- Attention pooling (2×2 mean-query) → beats simple feature stacking
- **Ordered points (top-down, left-right) → 86.3% vs 74.1% (unordered)**
- Plain-text coordinates → significantly better than special location tokens
- **Adding Indian-only document VQA data boosted performance on not only Patram-Bench, but also on other popular English DocVQA benchmarks**

Patram - India's First Document VLM

What is the IFSC Code in this Cheque

Indian Bank

Branch: TALOJA
SHOP 12 AND 3 AMAR HARMONY PLOT NO 22 SECTOR 4
TALOJA PHASE 1 NEAR PANCHAND
RAJWADI STATION NAVI MUMBAI 410208
IFSC Code: IDIB000T199

VALID FOR THREE MONTHS ONLY

Date: 01/06/2022

PAY Self Rohan Kumar

OR BEARER

या धारक को

RUPEES रुपये ० वरुण हजार रुपये 43/-

अदा करें ₹ 5000/-

CBS Code: 3028

Text

DEPENDRA A SHA

PLEASE SIGN ABOVE

987000030

PAYABLE AT PAR AT ALL OUR BRANCHES

BB2018 4000190831 287490* 31



PATRAM



Gemma 3

Correct Answer

IDIB000T199

Hallucination

3028

Document VLM Leaderboard: Patram Results

Model	Overall	DocVQA	Patram-Bench	VisualMRC
Qwen2.5-VL-7B-Instruct	0.8759	0.8722	0.6816	0.9169
Patram-7B-Instruct	0.8692	0.9373	0.6888	0.8598
Gemma-3-12B-IT	0.8556	0.8451	0.6349	0.9069
InternVL3-9B	0.7865	0.8681	0.6432	0.7405
DeepSeek-VL2	0.7581	0.8739	0.5089	0.7144
Molmo-7B-D-0924	0.6820	0.6314	0.5018	0.7577
DocOwl2	0.6776	0.7391	0.4569	0.5868
Molmo-7B-O-0924	0.6654	0.6314	0.5018	0.7577
Llava-v1.6-mistral-7b	0.2577	0.1263	0.4134	0.3889



DocEval
An Evaluation framework for DocAI Tasks



WASH Usecase: Compost Pit Pictures

Example 1: Image Blur + Lat Long extraction

Community Compost Pits (SBM)



Category :- Community Compost Pits



Lat:26.627175
Long:83.2842139

Viewing: community compost pits_image blurred_not clear.png

Clear View

Chat with Doc



Namaste Bharatgen! Welcome to your new DocBodh chat. How can I assist you?



is the image in the document blurred or clear?



The image in the document is blurred and abstract, making it difficult to identify specific details or elements.



can you give the lat and long?

[Attached: community compost pits_image blurred_not clear.png]



The latitude is 26.627175 and the longitude is 83.2842139.

WASH Usecase: Compost Pit Pictures

Example 2 - Image is compost pit or not + Lat long extraction



Lat:26.5506573
Long:83.358635

Viewing: community compost pits_wrong picture .png

Clear View

Chat with Doc



Namaste Bharatgen! Welcome to your new DocBodh chat. How can I assist you?



is this a picture of a compost pit inside the document?



The image does not show a compost pit. It depicts a natural, cluttered area with debris and vegetation.



what is the lat and long?



The latitude is 26.5506573 and the longitude is 83.358635.



A multilingual document VLM for India — 22+ languages, one model.



Multilingual by design

Trained natively.



Indian-context native

Code-mixed text, stamps, seals,
regional form layouts,
handwriting.



Linguistics in the loop

Makes the model correct, not just
fluent.