

A Low-Cost Screener for Myopia and Hypermetropia Using ResNet-18: Field-Capture Evaluation

Nelson R. Varte¹, Navajit Saikia², Saptaparna Bhattacharjee², Kristi Kayum Kuli², Shruti Paul²

¹Department of Computer Applications

²Department of Electronics and Telecommunication Engineering, Assam Engineering College, Guwahati 781013, India, All authors have equal contributions

Received: 1st August, 2026; **Revised:** 5th August, 2026; **Accepted:** 7th August, 2026; **Available Online:** 31st August, 2026

ABSTRACT

Uncorrected myopia and hypermetropia remain a common cause of avoidable visual impairment, and the shortfall is worst where trained optometrists and autorefractors are scarce. We ask a narrow, practical question: can a small convolutional network reading the Brückner red reflex run usefully on hardware a rural health worker already owns? We fine-tuned ResNet-18 on 203 clinically diagnosed red-reflex images (74 myopia, 122 hypermetropia, 7 normal) collected at a single eye hospital, and deployed it twice - first on a Raspberry Pi 5 serving a browser interface over a local network, then as a self-contained offline application on an ordinary Android phone, which removes the embedded board entirely. We evaluate on 300 images captured live through the phone under ophthalmoscope illumination, 100 per class, rather than on held-out curated images: with only 7 distinct normal-class eyes, a categorical held-out split of this dataset is too small to support a meaningful accuracy estimate, and we do not report one. On live captures the system reaches 67.3% accuracy (95% CI 61.8–72.4%), with per-class recall of 72% for myopia, 68% for hypermetropia and 62% for normal. End-to-end latency on the Pi is roughly one second per image. We report the live figure as the headline result because it is the one that reflects what the device does in a clinic corridor. Identifying these children matters because the main pharmacological option for slowing myopic progression, low-concentration atropine, is only useful when started during the years the eye is still elongating. The accuracy is not yet adequate for screening decisions; what the work establishes is that the deployment path is cheap, offline, and reproducible, and that dataset scale - not model capacity or hardware - is the binding constraint.

Index Terms - Myopia, hypermetropia, refractive error, Brückner reflex, ResNet-18, edge deployment, Raspberry Pi, on-device inference, vision screening, low-dose atropine, myopia control.

Source of support: Nil.

Conflict of interest: Nil

How to cite this article: Nelson R. Varte¹ Navajit Saikia² Saptaparna Bhattacharjee² Kristi Kayum Kuli² Shruti Paul². A Low-Cost Screener for Myopia and Hypermetropia Using ResNet-18: Field-Capture Evaluation. Int J Drug Deliv Technol. 2026;16(82s):892-900. DOI: 10.25258/ijddt.16.82s.98.

INTRODUCTION

Myopia and hypermetropia are among the most common visual impairments worldwide, and they are least well served in rural and low-income settings where ophthalmic care is thin on the ground. In a myopic eye an elongated axial length brings light to a focus in front of the retina, blurring distant objects. In hypermetropia a shorter axial length or a flatter cornea places the focus behind the retina, so



Catching these errors early matters because untreated refractive error in children can progress to amblyopia and strabismus, and in severe myopia to retinal detachment. Conventional screening - visual

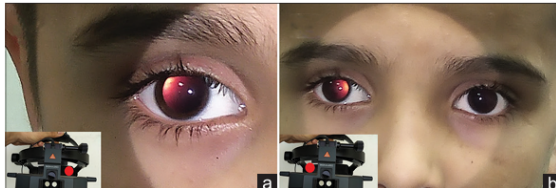
near work becomes difficult.

The Brückner test exploits a simple optical consequence of this. Shine a light into the pupil and the returning red reflex is not uniform: a myopic eye shows a darker red reflex across the upper pupil with a paler lower portion, a hypermetropic eye shows the reverse, and an emmetropic eye returns an even tone. Figure 1 shows the three patterns.

Figure 1. Red reflex in (i) myopia: upper pupil red, lower pale; (ii) hypermetropia: upper pale, lower red; (iii) normal: even tone.

acuity charts, retinoscopy, autorefraction - needs trained staff and equipment that a village clinic does not have. The Brückner test itself is cheap, which is why pediatricians have long performed it with a

direct ophthalmoscope (Fig. 2), but reading the reflex reliably is a skill. Reports of accelerated myopia progression in children following COVID-



This paper is about closing that skill gap with hardware that already exists in the field. Our interest is less in squeezing out another point of benchmark accuracy than in establishing whether the whole pipeline - capture, inference, result - can be made to run offline on a phone, and in measuring honestly what it achieves when it does.

A. Contributions

We built a clinically diagnosed red-reflex dataset in collaboration with an eye hospital, covering myopia, hypermetropia and normal vision.

We deployed a fine-tuned ResNet-18 on a Raspberry Pi 5 as a self-contained offline screening station reachable from any phone browser on the local network.

We then removed the embedded board altogether, running the classifier entirely on a stock Android phone - to our knowledge an unusually low floor for the hardware cost of an automated Brückner screener.



Under Brückner illumination the anterior focal point produces a darker, redder reflex over the upper pupil

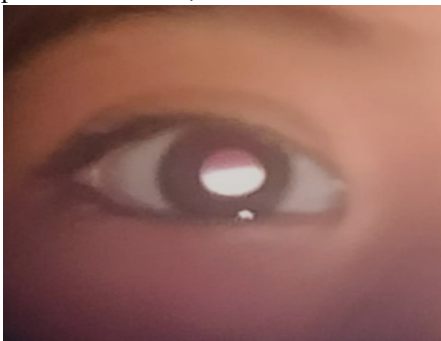


Figure 4. Myopic eye: red reflex on the upper pupil.

B. Hypermetropia

Hypermetropia arises from a short axial length, a flat cornea, or insufficient lens curvature, placing the focus behind the retina (Fig. 5). Children often mask mild hypermetropia through accommodation, which

19 home confinement [1], [2] have made the gap more pressing.

Figure 2. An eye captured for red-reflex observation.

We evaluate on live captures rather than curated images, and report the resulting 67.3% plainly, together with an analysis of why a dataset of this size cannot support the higher figures common in this literature.

II. Background and Related Work

A. Myopia

Myopia usually begins around school age, six to eight years, as the eye's axial length starts to outrun its optical power [3]. Progression during childhood and adolescence typically runs between -0.50 and -1.00 dioptres a year and generally stabilises in the late teens [3], [4]. Both inheritance and environment contribute; prolonged near work and reduced time outdoors are the environmental factors most consistently implicated. The scale of the problem is large: 80–90% of young adults are affected in parts of East Asia, and roughly half the world population is projected to be myopic by 2050 [4]. Figure 3 contrasts image formation in a normal and a myopic eye.

Figure 3. Image formation in a normal eye and a myopic eye.

and a brighter lower portion [6], [7], as in Fig. 4.

pupil.

is precisely why it escapes notice until it produces eye strain, headache or difficulty with close work. Prevalence rises with age, partly because presbyopia is folded into many published estimates.

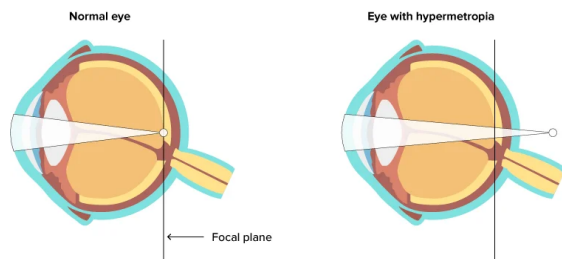


Figure 5. Image formation in a normal eye and a hypermetropic eye.

The Brückner pattern inverts: the red reflex sits on the lower pupil, as in Fig. 6.

C. Learning-based detection of refractive error

Classical machine learning has been applied to structured ophthalmic data - axial length, corneal curvature, anterior chamber depth, acuity scores - using support vector machines, random forests and gradient-boosted trees, generally reporting AUC above 0.75. Those methods need biometric measurements, which presupposes the equipment we are trying to do without.

Work closer to ours applies convolutional networks directly to smartphone photorefractometry or red-reflex images. Chun et al. [8] trained a custom residual network on smartphone photorefractometry images and reported 81.6% across seven dioptric classes. Linde et al. [9] estimated spherical power from red-reflex images captured with a low-cost handheld fundus camera at roughly 75%. Massmann et al. [10] reported about 90% for a smartphone red-eye reflex test, and Syauqie et al. [14] a multi-branch CNN reaching 91% with an AUC of 0.99. Da Silva et al. [15] took a different route, segmenting the reflex region itself with a U-Net.

These studies share two properties that ours does not: substantially larger datasets, and ground truth from cycloplegic refraction rather than categorical diagnosis. They also stop at the model. What none of them report is an end-to-end deployment measured on images captured through the deployed device itself, which is the gap this paper addresses - and, as our results show, the gap where most of the accuracy goes.

D. Screening as the entry point to pharmacological myopia control

Detection is not an end in itself, and it is worth being explicit about what a positive screen is for. The principal pharmacological option for slowing myopic progression in children is low-concentration topical atropine. The LAMP randomised controlled trial established a concentration-dependent effect across 0.05%, 0.025% and 0.01% atropine against placebo, with 0.05% the most effective of the three over one year [17], and five-year follow-up supports sustained treatment [18]. The evidence is not uniform across populations - 0.01% performed no better than placebo in a United States trial while showing benefit in East Asian cohorts - so the optimal concentration outside East Asia remains an open question.

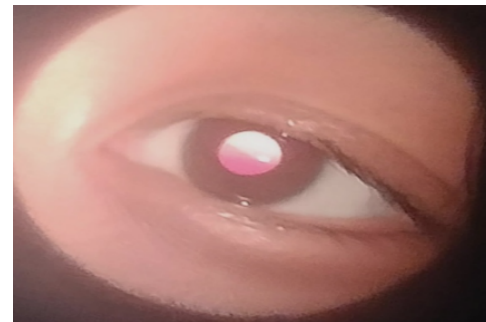


Figure 6. Hypermetropic eye: red reflex on the lower pupil.

What is not in question is that the therapy reaches only those children who have been identified. Atropine acts on progression rather than on established refractive error, so its value depends on beginning treatment while the eye is still elongating, broadly from six to eight years of age onward. A child whose myopia goes unnoticed until secondary school has spent most of that window untreated. Screening coverage is therefore an upstream constraint on the reach of pharmacological myopia control: no formulation, however well designed, reaches a patient nobody has found. In districts with one ophthalmologist per several hundred thousand people, that constraint binds harder than any property of the drug.

The formulation side compounds the same access problem from the other direction. There is no widely approved commercial low-dose atropine product, so the drug is typically supplied by compounding pharmacies, and surveys of United States compounders report wide variation in preparation method, diluent, storage and beyond-use dating, with consequences for potency and stability [19]. Stability characterisation of ophthalmic atropine solutions has become an active area in response; 0.1 mg/mL preparations with and without preservative have been shown to hold above 94% of initial concentration over six months in low-density polyethylene multi-dose droppers [20]. A rural screening programme in India faces both constraints simultaneously - securing a stable preparation, and finding the children who need it. This paper addresses only the second of the two and makes no claim regarding the first.

There is one further pharmacological dependency worth noting, because it shapes our dataset. The reference standard in most published refractive-error work is cycloplegic refraction, which requires instilling a cycloplegic agent and therefore a clinical setting and a waiting period. Our labels are categorical diagnoses rather than cycloplegic dioptric values. That is a genuine limitation for comparability, discussed in Section VII, but it is also what makes the screening step deployable outside a clinic.

E. Hardware

The direct ophthalmoscope (Fig. 7) supplies illumination. Its aperture disc selects apertures and

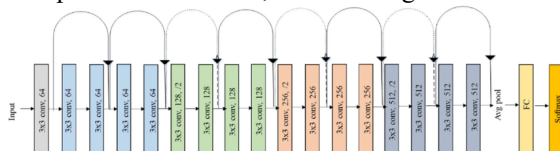
filters, and a lens wheel carrying roughly -20 D to



The Raspberry Pi 5 (Fig. 8) provides a quad-core Cortex-A76 at up to 2.4 GHz with up to 8 GB LPDDR4X, on a 32 GB Class 10 microSD card. We



ResNet-18 [11] uses residual blocks whose shortcut connections ease gradient flow through depth. Against ResNet-50 or ResNet-101 it is far lighter, which is what recommends it here: the trained model occupies about 45 MB, small enough to sit on a



III. Dataset

A. Acquisition

Publicly available red-reflex data labelled by refractive class is scarce, so we collected our own in collaboration with Sri Shankardeva Nethralaya Hospital. With institutional approval and under the applicable ethical guidelines, anonymised red-reflex images were obtained from patients already diagnosed by an ophthalmologist as myopic, hypermetropic or emmetropic. Images were included when the pupillary reflex was clearly visible, the clinical diagnosis was confirmed, and resolution was adequate; blurred, overexposed and underexposed frames were discarded.

B. Composition and its consequences

The resulting set is small and badly imbalanced: 74 myopia, 122 hypermetropia and 7 normal, totalling 203 images. The seven normal images are the central fact about this dataset and they constrain everything downstream. Seven eyes cannot represent healthy variation across age, iris colour, pupil size and **Table 1. Dataset before and after augmentation.**

+20 D lets the examiner bring the reflex into focus.

Figure 7. Direct ophthalmoscope used for red-reflex illumination.

used the 8 GB variant. It runs ResNet-18 comfortably, costs little, and draws little enough power to run from a battery.

Figure 8. Raspberry Pi 5 with a 32 GB Class 10 SD card.

phone alongside everything else. The network takes a 224×224 RGB input, applies a 7×7 convolution, then four residual groups of 64, 128, 256 and 512 filters, global average pooling to a 512-dimensional vector, and a fully connected classifier (Fig. 9).

Figure 9. ResNet-18 architecture.

refractive state, and no amount of augmentation changes that - augmentation redistributes existing information, it does not add any.

This also rules out the evaluation protocol one would normally reach for. A stratified 15% test partition of 203 images yields about 31 images, of which one would be a normal-class eye. A single image can only score 0% or 100%, so class-wise precision and recall for the normal class are undefined in any useful sense, and the overall accuracy moves in steps of 3.2 percentage points. We therefore do not report a curated held-out accuracy. Reporting one would attach a decimal point to a quantity the data cannot resolve.

Augmentation was applied to training data only, expanding it as shown in Table 1. Transformations were random rotation ($\pm 15^\circ$), horizontal flipping, brightness ($\pm 20\%$) and contrast ($\pm 15\%$) adjustment, Gaussian noise, zoom between $0.8\times$ and $1.2\times$, and random cropping with resizing [16]. Figures 10–12 show representative examples.

Class	Original images	After augmentation	Expansion
Myopia	74	1,006	13.6×
Hypermetropia	122	1,401	11.5×
Normal	7	107	15.3×
Total	203	2,514	12.4×

The expansion factors were chosen to even out class frequencies during training. Note what this means for the normal class: 107 training images derived from 7 eyes. The network sees each of those seven eyes about fifteen times in slightly different guises.

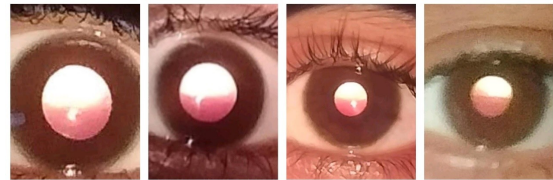


C. Organisation

Images were stored in a class-folder hierarchy (Myopia, Hypermetropia, Normal) so that PyTorch



Figure 13. Sample myopic eyes.



IV. Model and Training

We initialised ResNet-18 from ImageNet weights [5] and replaced the final fully connected layer with a three-output layer for myopia, hypermetropia and normal. Training used PyTorch and Torchvision [12], with the model on GPU where available and CPU otherwise. All images were resized to 224×224, converted to tensors and normalised with Table 2. Training progress. Validation accuracies are computed on the 383-image augmented validation partition.

Epoch	Train loss	Validation loss	Validation accuracy
1	0.2489	0.1867	92.95%

Figure 10. Augmentation: random rotation.



Figure 11. Augmentation: random cropping.



Figure 12. Augmentation: horizontal flipping.

and TensorFlow loaders infer labels from directory names. Representative samples of each class appear in Figures 13–15.

Figure 14. Sample hypermetropic eyes.

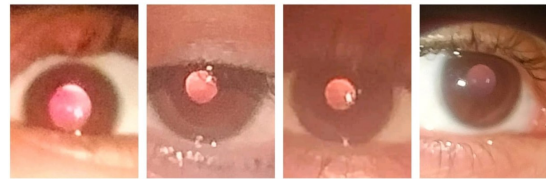


Figure 15. Sample normal eyes.

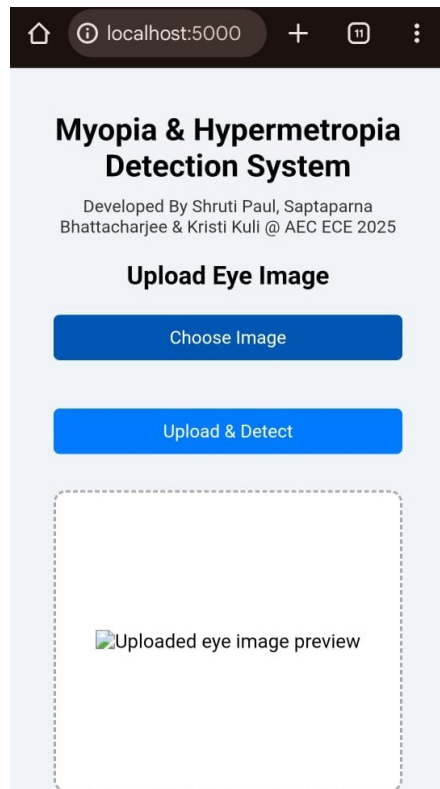
ImageNet channel statistics. A random horizontal flip was applied in the training pipeline only; validation used no random transformation, so that successive epochs are comparable.

Training ran for 10 epochs. Table 2 records the trend. Convergence was fast - training loss fell from 0.2489 to 0.0023, and validation accuracy exceeded 96% by the second epoch.

Epoch	Train loss	Validation loss	Validation accuracy
2	0.0385	0.1120	96.34%
3	0.0269	0.0623	97.91%

Epoch	Train loss	Validation loss	Validation accuracy
4	0.0135	0.0616	98.17%
5	0.0117	0.0301	99.22%
6	0.0160	0.1135	97.13%
7	0.0088	0.0560	98.17%

These numbers require a caveat that we want to state before anyone quotes them. The validation partition here is drawn from the augmented pool, so augmented variants of a given eye can appear on both sides of the split. Validation accuracy above 98% therefore measures the network's ability to recognise transformed copies of eyes it has already seen, not its ability to generalise to new patients. We report the trend because it documents that optimisation converged and that the model was not diverging; we do not report it as evidence of screening performance, and it should not be compared against the figures in Section VI or against published accuracies. Fluctuations in validation loss at epochs 6, 8 and 9 guided model selection.



Capture proceeds as in Figures 17–19: the eye is illuminated with the ophthalmoscope, the phone

Epoch	Train loss	Validation loss	Validation accuracy
8	0.0029	0.0919	98.17%
9	0.0019	0.1146	98.43%
10	0.0023	0.0712	98.69%

V. Deployment

A. Raspberry Pi 5

The trained model was first installed on a Raspberry Pi 5 running Raspberry Pi OS Lite with PyTorch 2.x, OpenCV 4.x and a Flask server. The Pi hosts a web page on the local network; a health worker points any phone browser at the Pi's address, uploads a captured image, and the classification comes back in the browser (Fig. 16). Nothing leaves the device and no internet connection is required, which is the point - connectivity is often the first thing missing at a rural camp. Throughput on the Pi is roughly one image per second, adequate for a workflow where a human is positioning an ophthalmoscope between captures.

Figure 16. Browser interface for upload and classification.

camera records the reflex, and the resulting frame is what the model receives.



Figure 17. Eye illuminated for reflex capture.



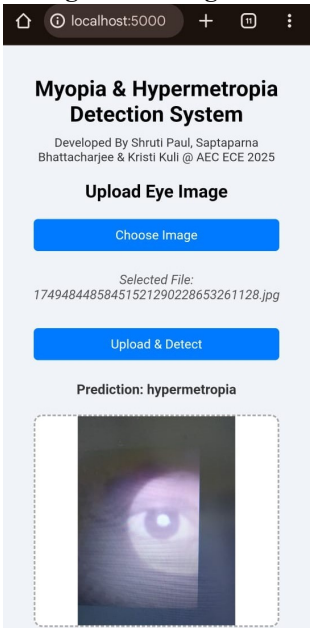
Figure 18. Capturing the Brückner reflex with a mobile phone.



Figure 19. Reflex as seen on the phone screen.

B. Removing the board: on-device inference

The Pi works, but it is still a second device to buy, carry and charge. We therefore exported the trained network to a portable interchange format [13], preserving the 224×224 input specification, and ran inference directly on an Android handset. The phone hosts the classifier locally and presents the same capture-and-result interface through its own



through this configuration. The marginal hardware cost of this configuration is an ophthalmoscope and a phone the user already owns. That is the practical contribution of the paper, and it is independent of how well the current model

performs. The approach has real drawbacks. Initial setup currently requires manual configuration beyond most non-technical users, the interface is browser-
browser; no embedded board, server or network connection is involved at any point. The complete on-device package is roughly 45 MB. Figure 20 shows a live capture classified on the handset. Implementation details of the on-device deployment are withheld pending an intellectual property filing. The evaluation in Section VI is unaffected: every image reported there was captured and classified
Figure 20. Live capture classified as hypermetropia on the phone.

performs. The approach has real drawbacks. Initial setup currently requires manual configuration beyond most non-technical users, the interface is browser-

based rather than a packaged application and cannot drive the camera directly, and the arrangement suits single-user field use rather than a multi-user clinic.

VI. Field-Capture Evaluation

We evaluated on 300 images captured live through the phone under ophthalmoscope illumination, 100 Table 3. Confusion matrix on 300 live captures.

Actual \ Predicted	Hyper.	Myopia	Normal	Total	Recall
Hypermetropia	68	18	14	100	68.0%
Myopia	15	72	13	100	72.0%
Normal	22	16	62	100	62.0%
Predicted total	105	106	89	300	-
Precision	64.8%	67.9%	69.7%	-	67.3%

Overall accuracy is 202/300, or 67.3%. A Wilson 95% confidence interval places the true value between 61.8% and 72.4%, so the figure should be read as approximate. Table 4 gives the per-class breakdown derived from the same matrix.

Table 4. Class-wise performance on live captures, derived from Table 3.

Class	Precision	Recall	F1	Support
Hypermetropia	64.8%	68.0%	66.3%	100
Myopia	67.9%	72.0%	69.9%	100
Normal	69.7%	62.0%	65.6%	100
Macro average	67.4%	67.3%	67.3%	300

Two patterns are worth drawing out. Myopia is recognised best, at 72% recall, which follows from it having the second-largest and most visually distinctive set of training examples. The normal class has the lowest recall at 62%, and 38 of its 100 captures were classified as pathological - 22 as hypermetropia, 16 as myopia. For a screening tool this is the costly direction of error: it generates false referrals, and a false referral rate near four in ten would overwhelm the specialist capacity such a programme is meant to conserve. We attribute it directly to the seven-eye training base for that class. Table 5 places the result alongside comparable studies, with an important caveat about what is being compared.

Table 5. Comparison with related red-reflex and photorefractive studies.

per class, with class membership given by the patient's clinical diagnosis. These are the conditions the device actually meets: handheld framing, ambient light varying with the room, occasional motion blur and corneal reflections. Table 3 gives the confusion matrix.

Study	Model	Task	Eval. data	Accuracy
Chun et al. [8]	Residual CNN	7-class dioptric	Curated	81.6%
Linde et al. [9]	CNN	Multiclass spherical	Curated	~75%
Massmann et al. [10]	CNN ensemble	Binary normal/abnormal	Curated	90% (F1 0.76)
Syauqi et al. [14]	3-branch CNN	4-class refractive	Curated	91%
This work	ResNet-18	3-class refractive	Live captures	67.3% (F1 0.67)

The comparison is not like-for-like and we would discourage reading the last row as a ranking. The four cited studies evaluate on curated image sets, with clinical refraction or expert annotation as ground truth and with substantially more data; we evaluate on images taken by hand through the deployed device. The tasks also differ. Massmann et al. classify normal against abnormal for visual impairments generally, on a dataset in which roughly seven images in ten are normal, so their 90% accuracy rests partly on that prevalence; their F1 of 0.76 against our macro F1 of 0.67 is the closer comparison and the gap is correspondingly smaller. What the table does establish is that the method family is sound, and that our dataset size rather than our architecture is what separates us from the cited results.

VII. Limitations

The dataset is the dominant limitation. 203 images from a single hospital cannot capture the demographic, device and illumination variation the system would meet in deployment, and the 7 normal-class eyes are too few to characterise healthy variation at all. Every result here should be read against that.

Labels are categorical rather than dioptric, so the model cannot express severity and cannot be compared directly against studies using cycloplegic refraction. Splitting is by image rather than by patient; with multiple images per patient this risks optimistic bias, and future collection should record patient identifiers to permit a patient-level split.

The live-capture set has 100 images per class, but we did not control the number of distinct individuals contributing to each class, so correlated captures may narrow the effective sample below 300. Capture is manual throughout, and framing, distance and ophthalmoscope alignment vary with the operator; some of the 32.7% error rate is attributable to capture technique rather than to the classifier, and we have not separated the two.

At 67.3% accuracy with 38% of healthy eyes flagged as pathological, the system is not usable for screening decisions today. We state this plainly because a tool aimed at children's vision should not be adopted on the strength of a promising deployment story.

VIII. Conclusion and Future Work

We set out to find whether automated Brückner screening could run on hardware a rural health worker already owns, and on that question the answer is yes: a ResNet-18 runs offline on a stock Android phone, with no embedded board, no connectivity and roughly 45 MB of storage. The Raspberry Pi 5 configuration works too and processes about one image per second. Measured on 300 live captures the classifier reaches 67.3% (95% CI 61.8–72.4%), which is not yet good enough to act on.

The gap is a data problem before it is a modelling

problem, so that is where the work should go next. The priority is a larger multi-centre dataset with a patient-level split, weighted heavily toward normal eyes - several hundred distinct healthy subjects would do more for this system than any change of architecture. Dioptric labels from cycloplegic refraction would allow severity estimation and direct comparison with the published literature. Beyond that, integrating a camera module directly with the Pi would remove the handheld phone and much of the framing variability with it, and a packaged native application would replace the manual setup that currently limits the system to technically confident users. Once accuracy is adequate, adding astigmatism, cataract and diabetic retinopathy would make the device a general screening tool rather than a two-condition one. The longer-term value of any such tool lies in what follows a positive result, so pairing screening coverage with a reliable supply of myopia-control therapy is the part of the problem that no classifier solves on its own.

Acknowledgment

The authors thank Sri Shankardeva Nethralaya Hospital for its support in data collection.

References

- [1] J. Wang, Y. Li, D. C. Musch, N. Wei, X. Qi, G. Ding, X. Li, J. Li, L. Song, Y. Zhang, Y. Ning, X. Zeng, N. Hua, S. Li, and X. Qian, "Progression of myopia in school-aged children after COVID-19 home confinement," *JAMA Ophthalmol.*, vol. 139, no. 3, pp. 293–300, 2021.
- [2] J. Wang et al., "Evaluation and follow-up of myopia prevalence among school-aged children subsequent to the COVID-19 home confinement in Feicheng, China," *JAMA Ophthalmol.*, vol. 141, no. 4, pp. 333–340, 2023.
- [3] I. G. Morgan, K. Ohno-Matsui, and S.-M. Saw, "Myopia," *The Lancet*, vol. 379, no. 9827, pp. 1739–1748, 2012.
- [4] B. A. Holden, T. R. Fricke, D. A. Wilson, M. Jong, K. S. Naidoo, P. Sankaridurg, T. Y. Wong, T. J. Naduvilath, and S. Resnikoff, "Global prevalence of myopia and high myopia and temporal trends from 2000 through 2050," *Ophthalmology*, vol. 123, no. 5, pp. 1036–1042, 2016.
- [5] J. Deng, W. Dong, R. Socher, L.-J. Li, K. Li, and L. Fei-Fei, "ImageNet: A large-scale hierarchical image database," in *Proc. IEEE CVPR*, 2009, pp. 248–255.
- [6] M. T. Kothari, "Can the Bruckner test be used as a rapid screening test to detect significant refractive errors in children?," *Indian J. Ophthalmol.*, vol. 55, no. 3, pp. 213–215, 2007.
- [7] P. Jain, M. T. Kothari, and V. Gode, "The opportunistic screening of refractive errors in school-going children by pediatrician using enhanced Bruckner test," *Indian J. Ophthalmol.*,

- vol. 64, no. 10, pp. 733–736, 2016.
- [8] J. Chun, Y. Kim, K. Y. Shin, S. H. Han, S. Y. Oh, T.-Y. Chung, K.-A. Park, and D. H. Lim, “Deep learning-based prediction of refractive error using photorefractive images captured by a smartphone: Model development and validation study,” *JMIR Med. Inform.*, vol. 8, no. 5, p. e16225, 2020.
- [9] G. Linde, R. Chalakkal, L. Zhou, J. L. Huang, B. O’Keeffe, D. Shah, S. Davidson, and S. C. Hong, “Automatic refractive error estimation using deep learning-based analysis of red reflex images,” *Diagnostics*, vol. 13, no. 17, art. 2810, 2023.
- [10] J. Massmann, A. Lichtenstein, and F. M. Lopez, “Early detection of visual impairments at home using a smartphone red-eye reflex test,” *arXiv preprint arXiv:2509.09808*, 2025.
- [11] K. He, X. Zhang, S. Ren, and J. Sun, “Deep residual learning for image recognition,” in *Proc. IEEE CVPR*, 2016, pp. 770–778.
- [12] A. Paszke et al., “PyTorch: An imperative style, high-performance deep learning library,” in *Proc. NeurIPS*, 2019, pp. 8024–8035.
- [13] ONNX, “Open Neural Network Exchange.” [Online]. Available: <https://onnx.ai>
- [14] M. Syauqie et al., “Deep learning for vision screening in resource-limited settings: Development of a multi-branch CNN for refractive error detection based on smartphone images,” *Frontiers in Computer Science*, vol. 7, art. 1576958, 2025.
- [15] I. F. S. da Silva, J. D. S. de Almeida, J. A. M. Teixeira, G. Braz Junior, and A. C. de Paiva, “Segmentation of the retinal reflex in Bruckner test images using U-Net convolutional network,” in *Image Analysis and Recognition (ICIAR 2018)*, LNCS vol. 10882, Springer, 2018, pp. 679–686.
- [16] C. Shorten and T. M. Khoshgoftaar, “A survey on image data augmentation for deep learning,” *Journal of Big Data*, vol. 6, art. 60, pp. 1–48, 2019.
- [17] J. C. Yam, Y. Jiang, S. M. Tang, A. K. P. Law, J. J. Chan, E. Wong, S. T. Ko, A. L. Young, C. C. Tham, L. J. Chen, and C. P. Pang, “Low-Concentration Atropine for Myopia Progression (LAMP) study: a randomized, double-blinded, placebo-controlled trial of 0.05%, 0.025%, and 0.01% atropine eye drops in myopia control,” *Ophthalmology*, vol. 126, no. 1, pp. 113–124, 2019.
- [18] X. J. Zhang, Y. Zhang, B. H. K. Yip, K. W. Kam, F. Tang, X. Ling, M. P. H. Ng, A. L. Young, P.-C. Wu, C. C. Tham, L. J. Chen, C. P. Pang, and J. C. Yam, “Five-year clinical trial of the Low-Concentration Atropine for Myopia Progression (LAMP) study: phase 4 report,” *Ophthalmology*, vol. 131, no. 9, pp. 1011–1020, 2024.
- [19] K. Richdale, E. S. Tomiyama, G. D. Novack, and M. A. Bullimore, “Compounding of low-concentration atropine for myopia control,” *Eye & Contact Lens*, vol. 48, pp. 489–492, 2022.
- [20] B. Berton, P. Chennell, M. Yessaad, Y. Bouattour, M. Jouannet, M. Wasiak, and V. Sautou, “Stability of ophthalmic atropine solutions for child myopia control,” *Pharmaceutics*, vol. 12, no. 8, art. 781, 2020.