



Vera C. Rubin Observatory
Systems Engineering

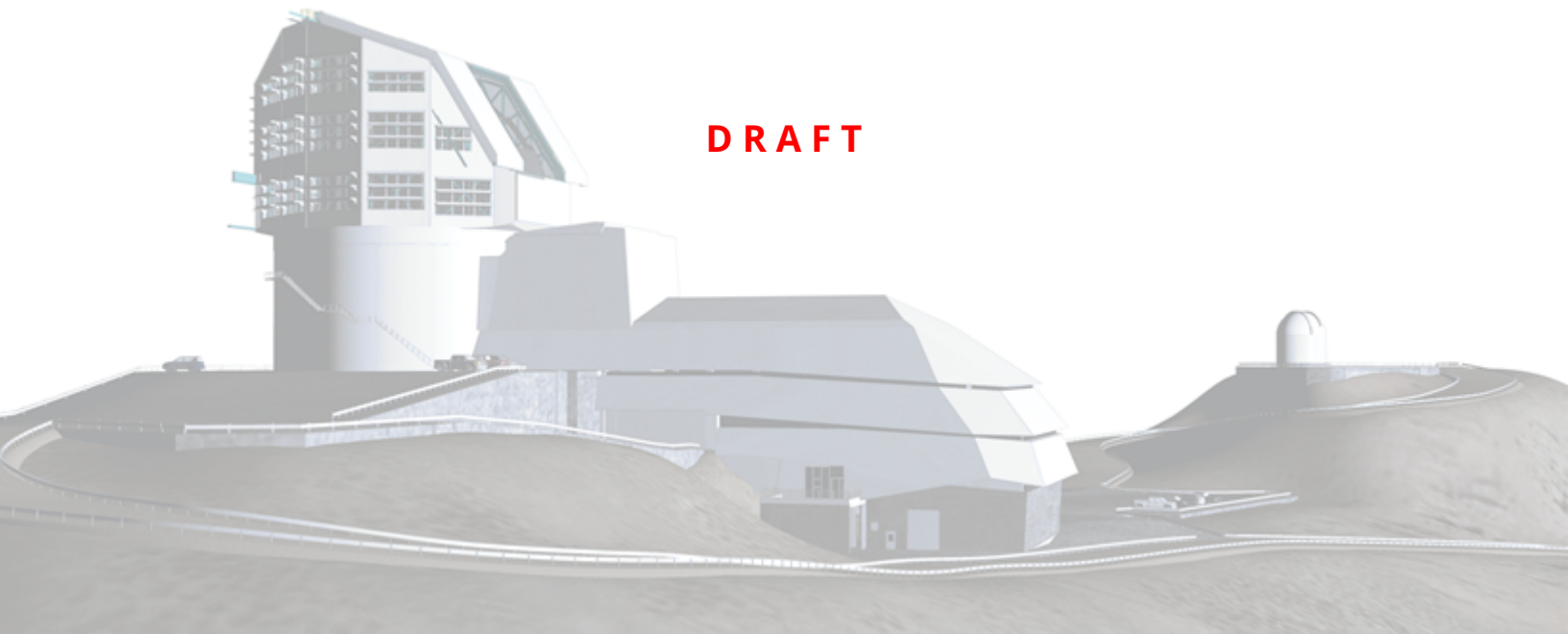
Initial studies of photometric redshifts with LSSTComCam from DP1

Eric Charles, John Franklin Crenshaw, Tianqing Zhang, Sam Schmidt,
Prakruth Adari, Johann Cohen-Tanugi, Melissa Graham, Julia
Gschwend, Bryce Kalmbach, Arun Kannawadi, Alex Malz, Sidney Mau,
Ignacio Sevilla-Noarbe, Rance Solomon, Dan Taranu

SITCOMTN-154

Latest Revision: 2025-07-09

DRAFT



Abstract

This technote holds reports based on the first analyses of the Data Preview 1 (DP1) data by the Science Unit for photometric redshifts. Although photometric redshifts are not an official DP1 data product, the “Photo-z Science Unit” generated photo-z estimates for every galaxy in DP1 using the available multi-band imaging on a best-effort basis. This work included developing training and test datasets by matching DP1 data to high-quality reference redshifts obtained with spectroscopy, Grism data, and multi-band photometry. The Science Unit used the RAIL software package to make photometric redshift estimates using eight different algorithms, developed simple scientific performance metrics, used those metrics to explore how the performance of the algorithms varied with configuration changes, derived more optimized configurations of the algorithms and tested the performance of those configurations. This work, the resulting data products and expected data distribution mechanism are all described there.

Change Record

Version	Date	Description	Owner name
1	2025-06-30	Initial release, coincident with release of DP1	Melissa Graham

Document source location: <https://github.com/lsst-sitcom/sitcomtn-154>

Draft

Contents

1	Introduction	1
2	Data	2
2.1	Rubin DP1	2
2.1.1	Preparation of photometric object catalogs	2
2.1.2	Data properties	4
2.2	Redshift reference sample	9
2.2.1	ECDFS redshift datasets	10
2.2.2	DESI Data Release 1 spectroscopic redshift dataset	11
3	Methodology	14
3.1	Template fitting based estimators	15
3.2	Machine learning based estimators	16
3.3	Analysis framework and bookkeeping software	16
3.4	Data exploration and algorithm optimization	16
3.5	Production of redshift catalogs	17
3.5.1	Catalogs in the rail_projects framework	17
3.5.2	Catalogs in the Rubin data management framework	17
4	Performance	18
4.1	Redshift estimator performance	18
4.1.1	Redshift estimation quality flags	18
4.1.2	Validation in the SV_38 field	21
4.2	Technical performance	21
5	Limitations and Caveats	24
6	Summary and Conclusion	25
A	Data products	28
A.1	Configuration files	28

A.2 Ancillary input files	28
A.3 Estimator data models	29
A.4 Redshift estimates stored as QP ensembles	29
A.5 Per-Object Point Estimates	30
A.6 Performance Monitoring Plots	32
B Data Distribution	32
B.1 Distribution via the Rubin Data Butler	33
B.2 Distribution via the Photo-z Server	33
B.2.1 From the PZ Server website	34
B.2.2 Via the pzserver Python library	34
B.3 Distribution as files at USDF and NERSC	35
B.4 Distribution via LSDB	35
C References	37
D Acronyms	39

Initial studies of photometric redshifts with LSSTComCam from DP1

1 Introduction

The Vera C. Rubin Observatory’s Data Preview 1 (DP1) marks an important milestone in preparing for the forthcoming Legacy Survey of Space and Time (LSST), offering a valuable opportunity to test and validate scientific tools and workflows on precursor imaging data (Vera C. Rubin Observatory, RTN-095). Among the core scientific objectives of LSST is the estimation of photometric redshifts (photo- z s) for billions of galaxies, enabling extragalactic astrophysics and cosmological analyses that rely on redshift estimates and distributions. Accordingly, the Rubin project developed a roadmap to providing high-quality photo- z s for the scientific community (Graham et al., DMTN-049).

Although photo- z are not a DP1 deliverable, the “Photo- z Science Unit” was asked to generate photo- z estimates for every galaxy in DP1 using the available multi-band imaging on a best-effort basis, laying the groundwork for future large-scale applications. This effort required integrating realistic data processing with scalable machine learning techniques capable of delivering precise redshift predictions across varied galaxy populations.

To accomplish this task, we employed the RAIL (Redshift Assessment Infrastructure Layers) software package, a flexible and modular platform designed for photo- z estimation and evaluation (The RAIL Team et al., 2025). Specifically, we used RAIL to train photo- z estimation algorithms on high-quality redshift training sets cross-matched to the DP1 photometric catalog.

This note describes the resulting best-effort photo- z catalogs, which should be regarded as extremely preliminary, and the tools used and the steps taken to generate those catalogs. In addition, we describe relevant features of the DP1 photometric data, the spectroscopic calibration datasets and the resulting photo- z catalogs. Additional details about data products and data distribution are included as appendices.

We expect feedback from science users as they explore the data and discover issues that we have not anticipated, and that we will incorporate that feedback in future work.

2 Data

2.1 Rubin DP1

The Rubin Observatory’s Data Preview 1 (DP1) dataset is the first public release of Rubin observatory imaging data processed through the LSST Science Pipelines (Rubin Observatory Science Pipelines Developers, PSTN-019), serving as a critical testbed for scientific and technical validation ahead of full LSST operations. DP1 is based on observations from the LSST commissioning camera (LSSTComCam) and includes multi-band optical imaging in u, g, r, i, z and y filters over several square degrees of sky. The dataset consists of processed images, source catalogs, and associated metadata, all formatted using the Rubin Data Butler system (Jenness et al., 2022) in the same way as full LSST data products. Although smaller in scale than future LSST datasets, DP1 offers realistic photometric measurements, object detection, and data structures, making it an invaluable resource for developing and testing algorithms for tasks such as photo-z estimation, object classification, and data quality assessment.

Critically, the fields selected for ComCam observation included the Extended Chandra Deep Field South (ECDFS), for which many high-quality redshifts exist from other surveys.

2.1.1 Preparation of photometric object catalogs

The Rubin Data Management (Rubin DM) pipeline measures multiple types of object photometry, the first stage in creating a photo-z catalog was to determine which measured photometry we would use as inputs. In this note, we generally use the 1.0 arc-second Gaussian Aperture fluxes and their associated errors, e. g. `u_gaap1p0Flux` and `u_gaap1p0FluxErr`. These fluxes should provide good measures of consistent galaxy colors within the defined aperture, ideal for photo-z estimation, though they may not necessarily reflect the colors of the overall galaxy if there is a significant color gradient and the galaxy is larger than the 1.0 arc-second aperture. This choice of photometric measurement may not be ideal, and investigation into the optimized set of photometric inputs will continue into the future. The only exceptions to this are that we used `i_psfFlux` in the initial broad data cuts, and that we briefly explored several of the other fluxes as part of the initial validation of our photo-z analysis pipelines.

Preparing object catalog (NSF-DOE Vera C. Rubin Observatory, 2025) data for photo-z algorithm training and estimation included five steps:

1. Applying quality cuts to the object catalog. We developed selections for the training and test datasets (*match_XXX*), and two additional selections for full DP1 catalogs: one applicable for fields with observations in only four bands (*gold_4_band*), the other for fields with observations in all six bands (*gold*). These selections are described below.
2. Converting fluxes in nJy to AB magnitudes ($m_{AB} = -2.5 \log_{10}(f_v/nJy) + 31.4$)
3. De-reddening to account for Galactic dust. We use the SFD dust maps (Schlafly & Finkbeiner, 2011).
4. Cross-matching objects with reference catalogs that include redshift information as described in Sec. 2.2.
5. Shuffling and splitting the resulting catalog into “training” and “test” data sets.

The data selection criteria that we used were:

- *match_prelim*: first version of matching to high-quality redshift catalogs in the ECDFS field, see Sec. 2.2.
- *match_ecdfs*: updated matching to high-quality redshift catalogs in the ECDFS field, see Sec. 2.2.
- *match_desi*: matching to DESI redshift catalogs in the **Rubin SV 38 7** field, see Sec. 2.2.2.
- *gold*: Detection in 'ugrizy' && $i_psfFlux / i_psfFluxErr > 5$ && ($g_extendedness > 0.5$ || $r_extendedness > 0.5$).
- *gold_4_band*: Detection in 'griz' && $i_psfFlux / i_psfFluxErr > 5$ && ($g_extendedness > 0.5$ || $r_extendedness > 0.5$).

To create the “training” and “test” data sets, we followed steps 1–5, while for the larger unlabeled data sets we followed steps 1–3. Specifically, for DP1, the **ECDFS**, “EDFS” and “Rubin SV 95 -25” fields have observation in 6 bands, and we applied the *gold* target selection. In **Rubin SV 38 7**, DP1 only has observation in 4 bands (“griz”), therefore we applied the *gold_4_band* target selection to that field.

All of these datasets are summarized in Tab. 1.

Data set	Selection	Number of objects
DP1	None	2,299,757
ECDfs+EDfs+SV_95	<i>gold</i>	375,610
SV_38	<i>gold_4_band</i>	169,034
training_v1	<i>match_prelim</i>	7,000
test_v1	<i>match_prelim</i>	2,437
training_v2	<i>match_ecdfs</i>	4,803
test_v2	<i>match_ecdfs</i>	1,201
test_DESI	<i>match_desi</i>	2,728

TABLE 1: Summary of the datasets used in this work. Note that to train models on four-band photometry, we used the same training and test sets as for six-band photometry, but configured the algorithms only to use the ‘griz’ bands.

2.1.2 Data properties

We have developed tools to generate diagnostic plots of both the input object catalogs and the photo- z estimates as part of our data analysis. Fig. 1 shows histograms of the AB magnitudes in 1.0 arc-second apertures of objects in the “test_v1” dataset, which includes an explicit magnitude cut, $m_i < 26.0$. Fig. 2 shows the correlation between magnitude and redshift for all objects in the “test_v1” data set. Fig. 3 shows the “adjacent band colors”, i.e., $u - g$, $g - r$, $r - i$, $i - z$, $z - y$ versus redshift for the same, with a series of SED templates (as used by template-fitting algorithms, see Sec. 3.1) overlaid. Finally, Fig. 4 shows the color-color plots for the same adjacent colors. In all cases, the 1.0 arc-second aperture magnitudes are used for plotting.

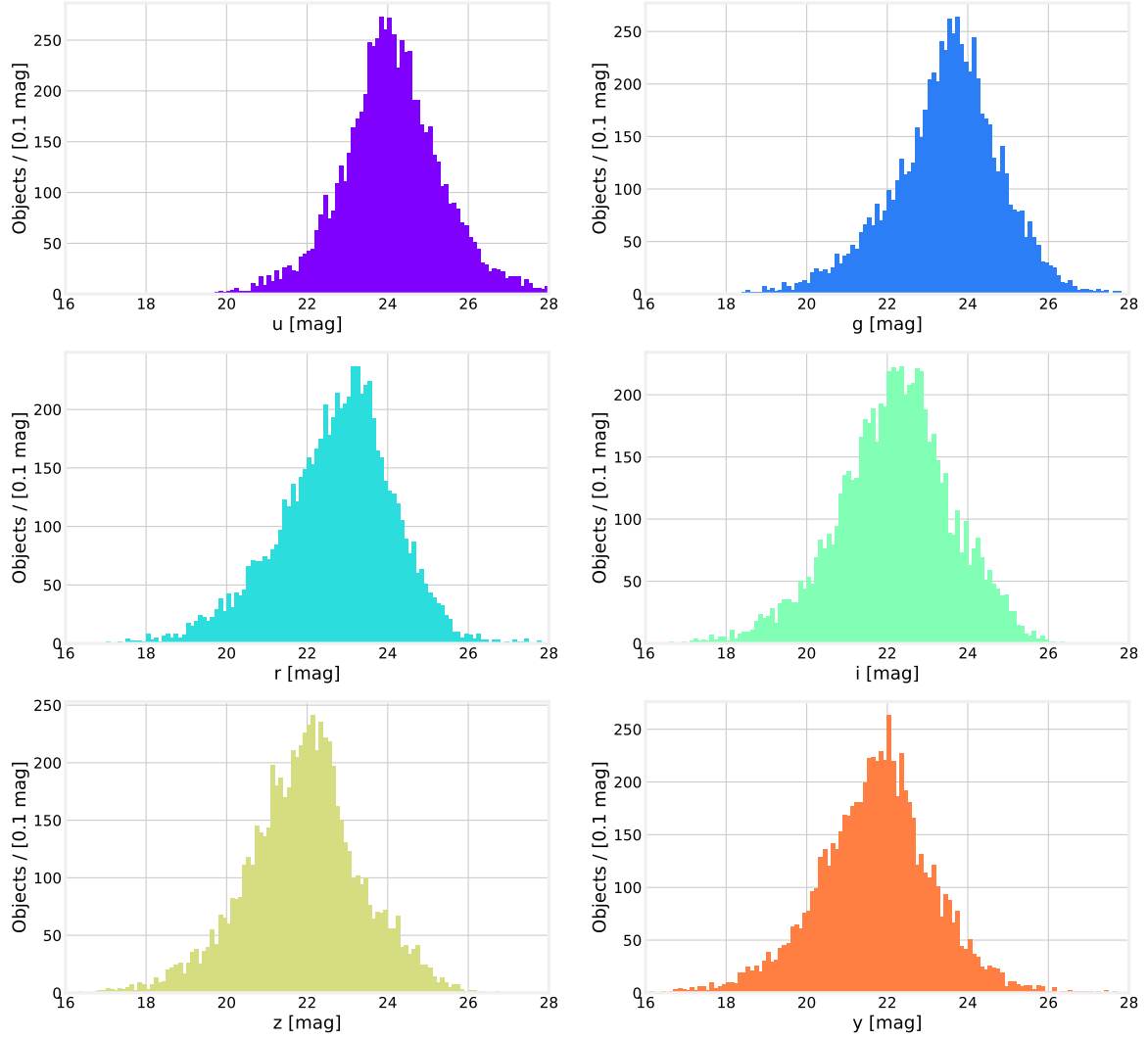


FIGURE 1: 1.0 arc-second Gaussian Aperture (Gaap) magnitudes (in AB system) of objects in each of the six Rubin filter bands in the “test_v1” dataset. These aperture magnitudes were used as the inputs for the photo-z algorithms.

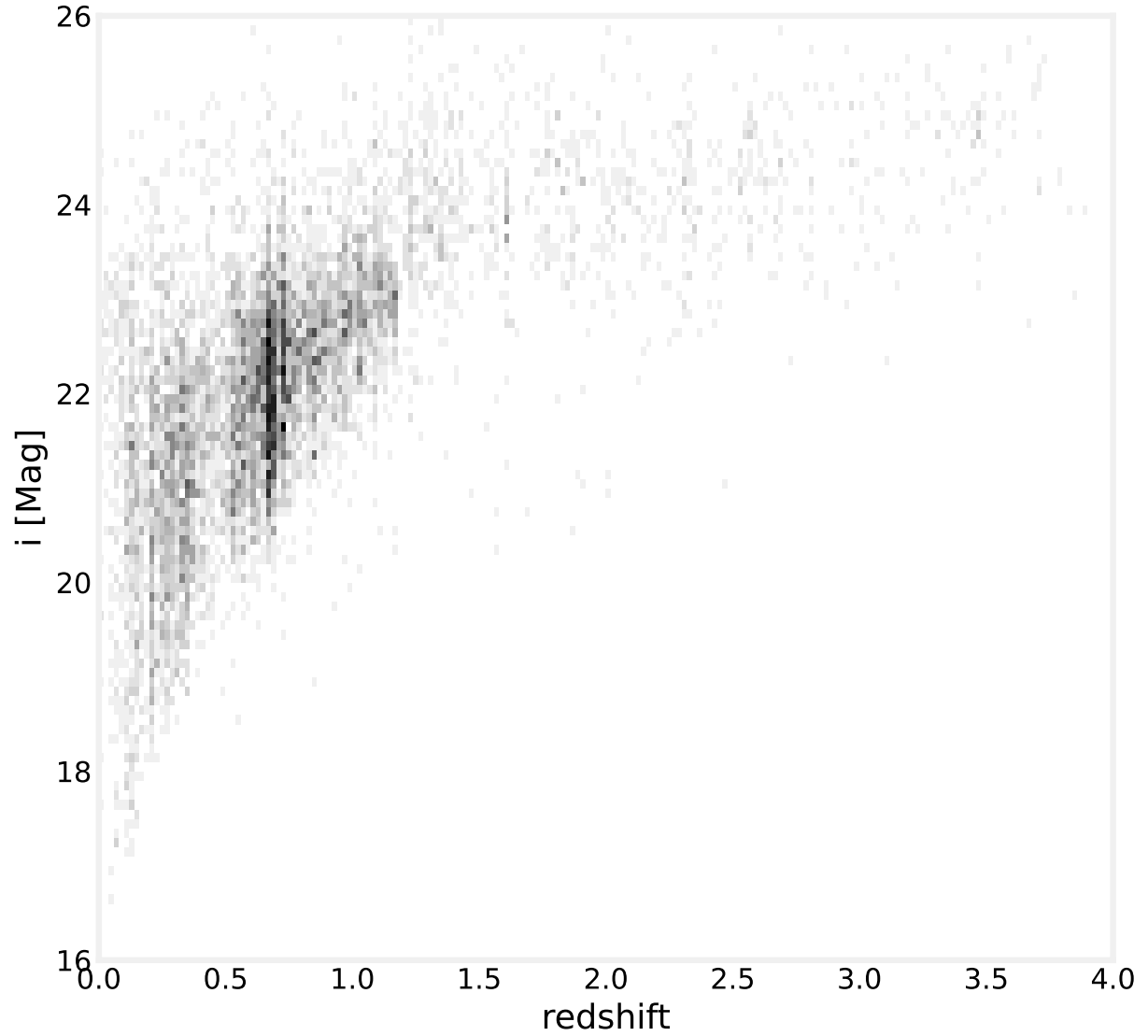


FIGURE 2: 1.0 arc-second Gaussian Aperture (Gaap) i-band magnitude versus redshift for all objects in the “test_v1” dataset.

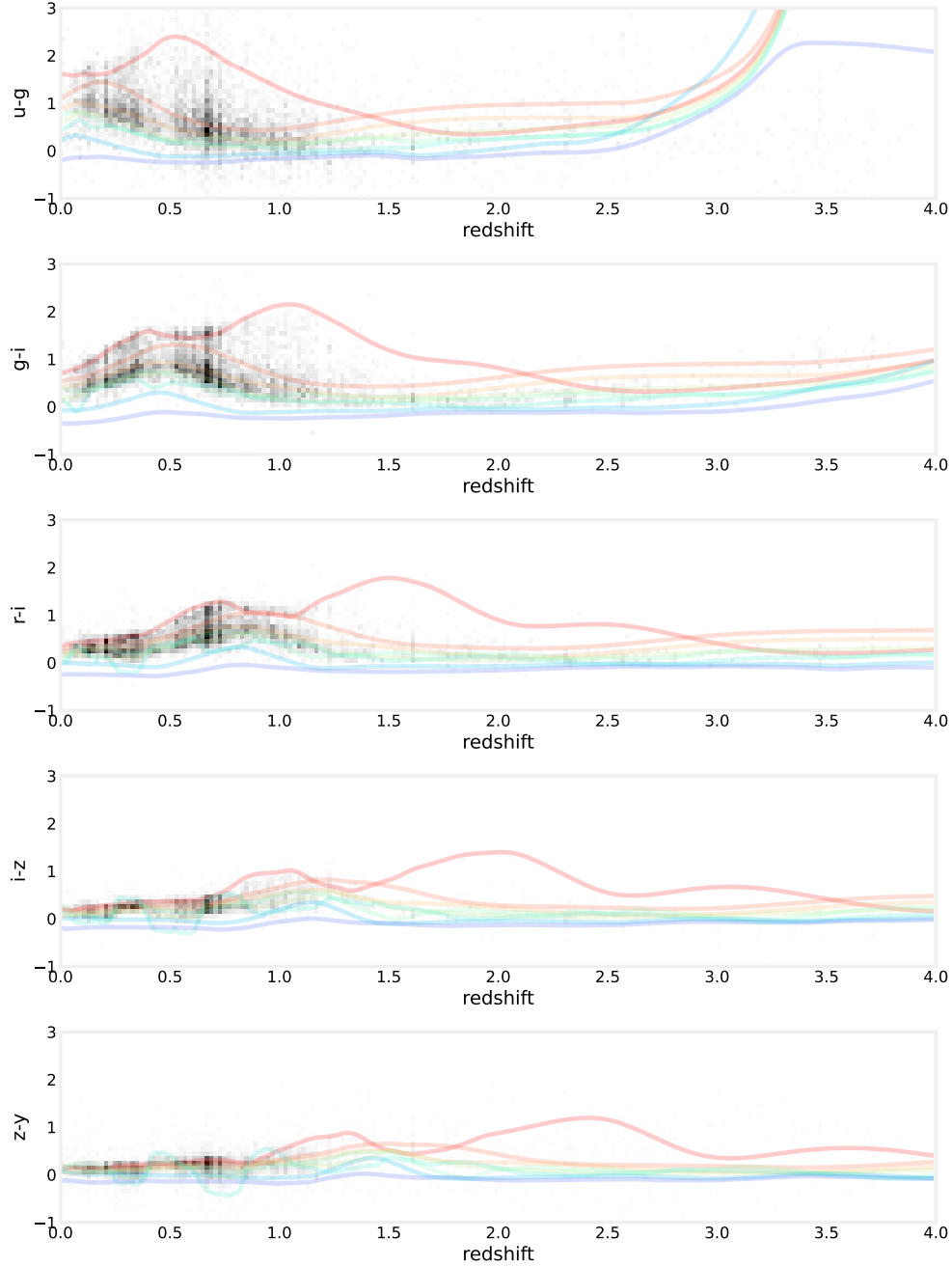


FIGURE 3: “Adjacent band colors”, i.e., $u - g$, $g - r$, $r - i$, $i - z$, $z - y$, in 1.0 arc-second apertures versus redshift for all objects in the “test_v1” dataset. Colored lines represent the expected colors for the eight “CWWSB” SEDs described in Sec. 3.1, and should very roughly show the predicted range of color evolution expected for our galaxy sample.

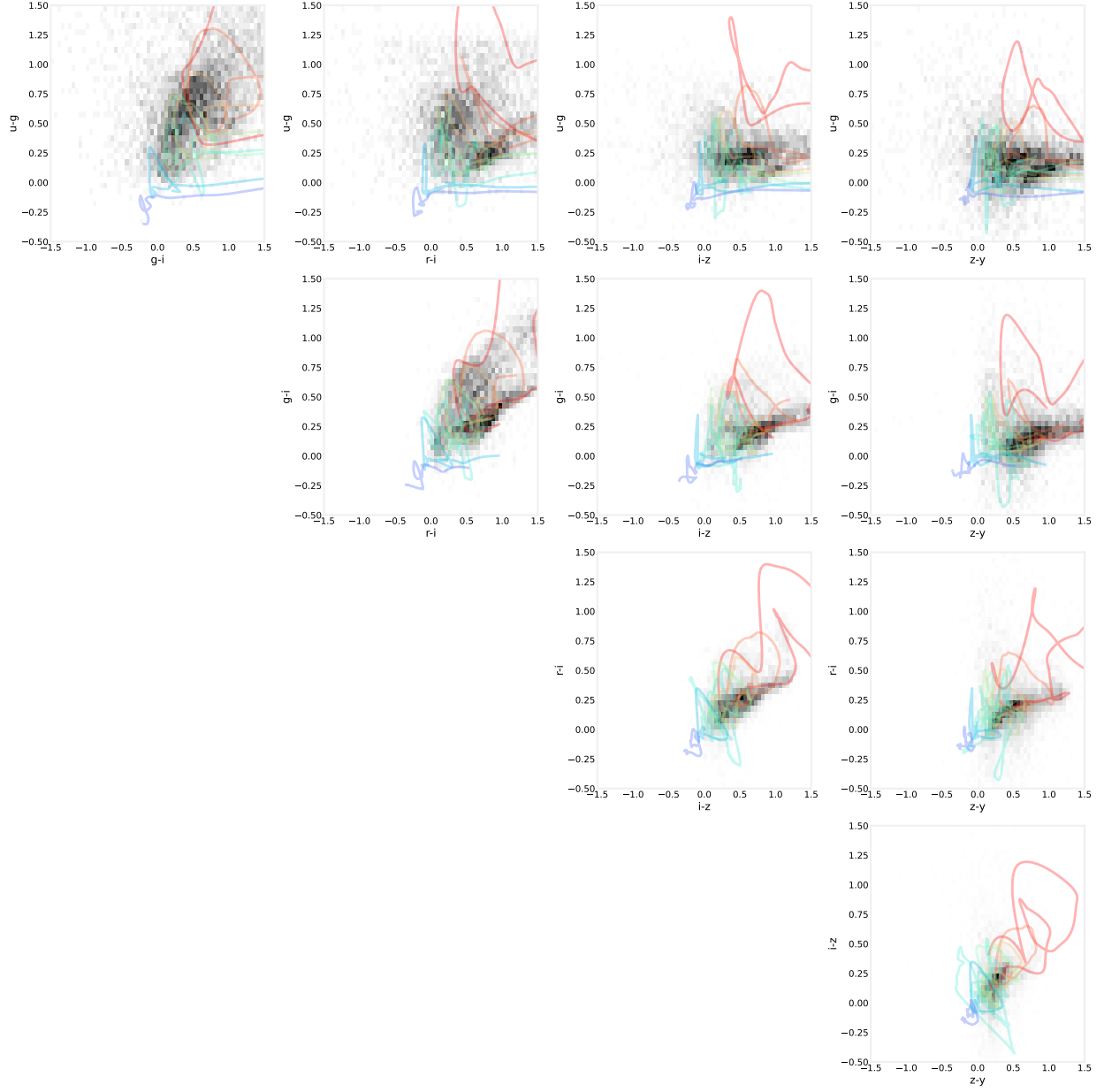


FIGURE 4: Color-color plots for all objects in the “test_v1” dataset. Colored lines represent the expected colors for the eight “CWWSB” SEDs described in Sec. 3.1, and should very roughly show the predicted color evolution expected for our galaxy sample.

2.2 Redshift reference sample

The photo- z estimation algorithms in RAIL require highly accurate and precise redshift estimates cross-matched to the DP1 photometric catalog to provide labeled “training” data sets. We created such datasets using data drawn from publicly available catalogs in the ECDFS, comprising galaxies with known spectroscopic redshifts, grism redshifts, and high-quality photo- z ’s from deep multi-band imaging. These training sets enable machine learning models within RAIL to learn the mapping between galaxy colors and redshifts, enabling photo- z estimation for every galaxy detected in DP1.

The reference data set used to train photo- z estimation algorithms can have an outsize impact on the resulting photo- z estimates, particularly for machine learning based methods. In many ways, the galaxies with known redshifts define the flux/color to distance relation by tracing out the mapping from empirical magnitudes to redshift that the algorithms “learn”. As such, the details of the construction of the reference sample is very important. In an ideal case, we would prefer to have a “representative” sample of redshifts, i. e. a fair sampling in terms of the color and magnitude distribution for all galaxies; however, due to the practicalities of spectrographs and expensive telescope time investment needed for deep spectroscopic campaigns, we must deal with incomplete training samples, particularly for faint and high redshift objects. We must also determine which datasets contain “secure” redshifts that meet some confidence threshold, and whether to include datasets from grism and many-band photo- z estimates that may contain a small fraction of incorrect redshift identifications. We continue to refine our reference sample definition, and we may include or exclude additional samples

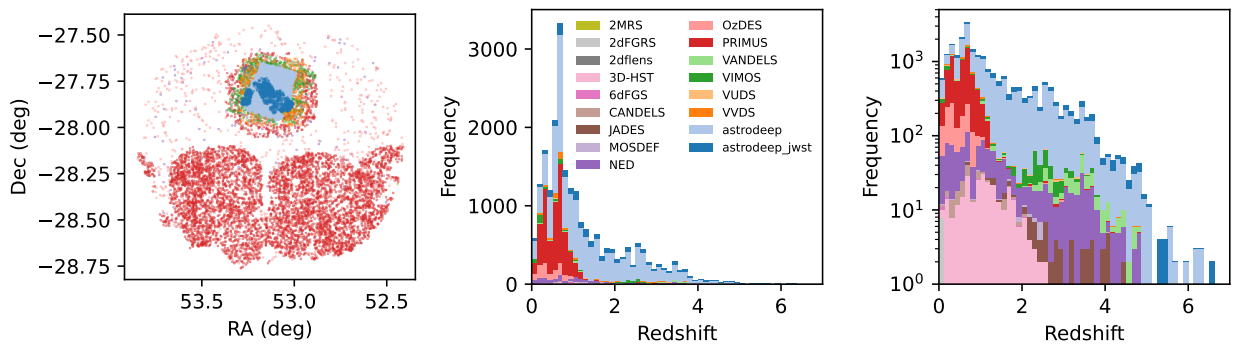


FIGURE 5: Redshift reference sample in ECDFS. Left: Locations of DP1 galaxies cross-matched to objects in each reference catalog (using color scheme from middle panel). Many of the sets overlap in the densely-covered GOODS-S field. Middle: Redshift distribution of objects in the reference catalogs. Right: Same as middle panel with a log scale on the y-axis.

as we investigate system performance, this note represents a current best effort, but the final selection will likely evolve. There are also some concerns of imprinted sample variance, as the deep six-band data of DP1 is concentrated in the single ECDFS field. Future Rubin data will cover multiple widely separated deep fields containing rich spectroscopic datasets, which will mitigate these concerns, but they may be an issue for the current DP1 estimates.

2.2.1 ECDFS redshift datasets

We have compiled and used two redshift samples in the ECDFS, (a preliminary version and a reference sample), consisting of spectroscopic-, grism-, and high-quality multiband photometric-redshifts (Fig. 5).

This reference sample is used to train machine learning photo- z estimators and evaluate photo- z performance. The component redshift catalogs (described below) were combined into a single reference catalog and their respective quality flags were homogenized by defining a redshift “confidence” (more on this below).

When combining the component redshift catalogs, sources within $0.75''$ were identified as duplicates. For these sources only the highest quality redshift is kept, i.e. spectroscopic redshifts are preferred over grism redshifts, which are preferred over photo- z ’s, and higher confidence values are preferred for redshifts of the same type. The redshift reference catalog was then cross-matched to the ComCam DP1 catalog using a radius of $0.75''$.

Confidence, which takes values between 0.0 and 1.0, is loosely defined as the fractional probability that an individual redshift estimate is correct. Most of the spectroscopic sets provide these estimates for their redshifts. For the few that don’t we assigned the confidence 0.95. For the grism and multiband photo- z surveys, we set the confidence equal to $1 - f_{\text{out}}$, where f_{out} is the reported outlier rate of these catalogs. To facilitate custom quality cuts, the catalog contains flags indicating whether each redshift originates from spectroscopy (`type == 's'`), grism (`'g'`), or multiband photo- z (`'p'`), as well as confidence values. Note redshifts from grism and photo- z surveys, however, have larger scatter and bias than spectroscopic surveys (and possible incorrect redshifts if the redshift is based off of a single emission line), but these metrics are not captured by the confidence parameter. If doing their own studies with custom reference samples, we encourage readers to investigate the details of each component survey that comprise the reference catalog and apply their own quality cuts as suit their needs. For example, if studying high-redshift galaxies, you might choose to include the multiband

photo- z 's, which increase the number of redshifts at $z > 2$ by a factor of 3 (Fig. 6).

We applied conservative cuts for our fiducial analyses, specifically type == 's', confidence ≥ 0.95 , SNR ≥ 10 in the i band (using `gaap1p0` fluxes). We then performed a random 80%/20% train/test split on this catalog, resulting in a training set of 4803 redshifts and a test set of 1201 redshifts.

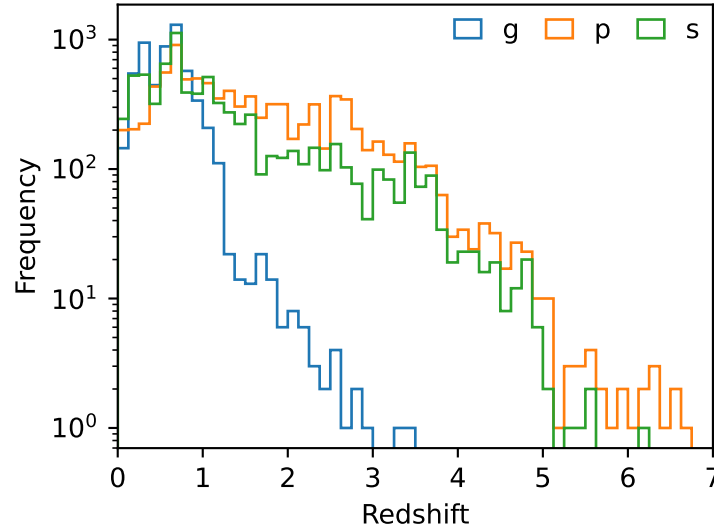


FIGURE 6: Redshift distribution of the reference catalog by redshift type, denoted s = spectroscopic, g = grism, p = multiband photo- z .

The *match_prelim* selection used in the “train_v1” and “test_v1” included all spectroscopic and grism galaxies in the “match_ecfds” catalog, without applying any confidence cut, and only required SNR ≥ 5 in the i band (using `psf` fluxes). This selection also dropped all galaxies with NaN magnitude in any of the six bands. We randomly selected 7,000 objects for the “train_v1” and assigned the remaining 2,438 to the “test_v1” dataset. Although we used these datasets extensively in V1 of this technote, we intend to update the note to use the “train_v2” and “test_v2” as soon as possible, ideally in the first week of July, 2025.

2.2.2 DESI Data Release 1 spectroscopic redshift dataset

We also use data from the DESI Data Release (Collaboration et al., 2025) to act as an independent validation data set as it overlaps the **Rubin SV 38-7 (SV_38)** field. The observations of this field are centered on the Abell 360 galaxy cluster which is located at $z = 0.22$ (Quintana & Ramirez, 1995). We cross-matched the DESI Bright Galaxy Sample (BGS) (Hahn et al., 2023),

Survey	Type	Confidence	Matches	Reference
2dFGRS	s	1.00	3	Colless et al. (2001)
		0.99	4	
		0.90*	1	
2df lens	s	1.00	1	Blake et al. (2016)
2MRS	s	0.95	1	Huchra et al. (2012)
6dFGRS	s	0.98	2	Jones et al. (2009)
3D-HST	g	0.99	5	Momcheva et al. (2016)
		0.95	277	
ASTRODEEP	s	1.00	4165	Merlin et al. (2021)
	p	0.97	8212	
ASTRODEEP-JWST	s	1.00	594	Merlin et al. (2024)
	p	0.92*	628	
		0.90*	455	
CANDELS	s	1.00	53	Kodra et al. (2023)
	p	0.93*	6	
JADES	s	0.99	11	D'Eugenio et al. (2025)
		0.95	34	
		0.90*	24	
MOSDEF	s	0.99	9	Kriek et al. (2015)
NED	s	0.95	847	Helou et al. (1991)
OzDES	s	0.99	897	Lidman et al. (2020)
PRIMUS	g	0.92*	3653	Cool et al. (2013)
		0.85*	1687	
VANDELS	s	1.00	196	Garilli et al. (2021)
VIMOS	s	1.00	499	Balestra et al. (2010)
		0.95	43	
VUDS	s	1.00	9	Tasca et al. (2017)
		0.95	9	
		0.80*	3	
VDS	s	1.00	101	Le Fèvre et al. (2005)
		0.95	193	
Total	s		7699	
	g		5622	
	p		9301	
	all		22622	

TABLE 2: Component surveys of the redshift reference sample. Redshift type is denoted s = spectroscopic, g = grism, p = multiband photo-z. Note for our fiducial analyses we applied conservative cuts on this catalog. Specifically, type == 's', confidence ≥ 0.95 , SNR ≥ 10 in the i band (using gaap1p0 fluxes).

* Note: multiband photo-z redshifts and grism and spectroscopic redshifts with confidence < 0.95 were not used in any training sets employed in this note.

Emission Line Galaxy (ELG) (Raichoor et al., 2023), and Luminous Red Galaxy (LRG) (Zhou et al., 2023) samples against the DP1 catalog using a radius of 0.5" producing 2728 matches with 398 from the BGS sample, 1421 from ELG, and 909 from LRG. The area of overlap and the three subsamples are shown in Fig. 7. This spans a redshift range of 0 to 1.6 and i -mag of 14 to 23.9 as shown in Fig. 8; the bump at $z \approx 0.25$ can be attributed to the cluster at the center of the field.

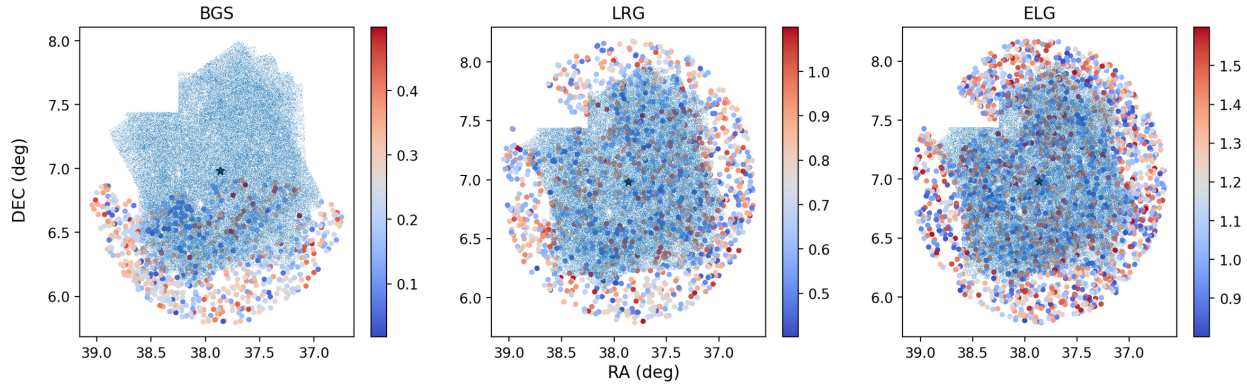


FIGURE 7: Overlap between the **SV_38** sample and DESI subsamples. In all plots, the small blue dots are DP1 objects from LSSTComCam and we overlay, from left to right, the BGS, LRG, and ELG subsamples with points color-coded by redshift. Each sample probes a different redshift range and is shown with the color bar to the right of each panel.

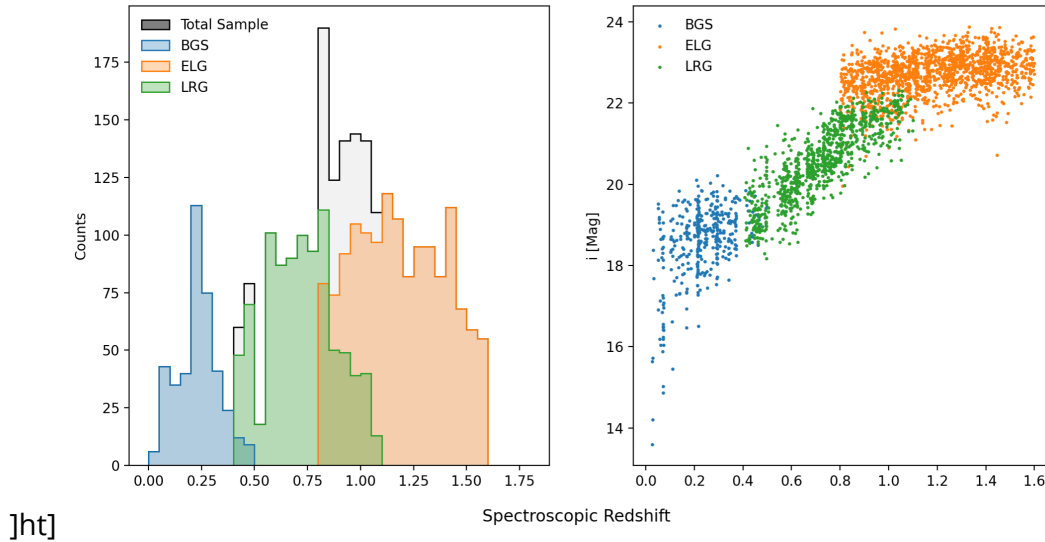


FIGURE 8: Validation set on **SV_38** from DESI BRG, ELG, and LRG samples. Left plot: the redshift distribution of the matched subsample with BGS in blue, ELG in orange, LRG in green, and the total matched distribution in black. Right plot: i -band magnitude as measured by LSSTComCam versus spectroscopic redshift for all DESI matches with subsamples shown using the same color scheme as left plot.

If an object was included in multiple samples, when matching we restrict to the one with the largest weight as calculated by DESI DR1. The **SV_38** field was observed in g (44 visits), r (55 visits), i (57 visits) and z (27 visits) bands restricting our validation to the 4 band models. Furthermore, the incomplete overlap with the BGS sample limits our ability to validate on redshifts below $z = 0.5$.

3 Methodology

We used the RAIL as the core tool for training and applying photo- z estimation models. We also used RAIL to evaluate the model performance through a suite of diagnostic metrics using photo- z point estimates, including redshift bias, scatter (e.g., normalized median absolute deviation), and catastrophic outlier rate and to make diagnostic plots of the algorithm’s performance. Specifically, we used the methodology described below to evaluate all the algorithms listed in Tab. 3.

As part of the Rubin photo- z roadmap process, KNN and BPZ were chosen as “testing algorithms” from the shortlist of algorithms considered in (Graham et al., DMTN-049), the expectation is that these will be supported by project data management team for Data Preview 2, while the others will be supported by the wider community, the RAIL development team and DESC collaboration.

Algorithm name	Home package	Reference
Project Supported Testing		
BPZ	rail-bpz	Benítez (2000)
KNN	rail-sklearn	RAIL Paper
Community Supported		
CMNN	rail-cmnn	Graham et al. (2018)
DNF	rail-dnf	De Vicente et al. (2016)
FlexZBoost	rail-flexzboost	Izbicki & Lee (2017)
GPz	rail-gpz-v1	Almosallam et al. (2016)
LePHARE	rail-lephare	Arnouts et al. (1999)
TPZ	rail-tpz	Carrasco Kind & Brunner (2013)

TABLE 3: Summary of the pre-wrapped estimators/summarizers/classifiers used in this paper and described detail in (The RAIL Team et al., 2025).

For each algorithm, we trained photo- z models using the reference datasets described in Sec. 2.2. Specifically, we used the `rail_project` package (see 3.3 to run RAIL’s PzPipeline. The pipeline consists of “Informing” or training the models on the “training” dataset, and using

those models in “Estimation” and “Evaluation” stages on the reserved “test” dataset to evaluate the algorithm performance. We also used `rail_project` to run RAIL’s EstimatePipeline, to perform the photometric redshift estimation on the unlabeled data sets for all of the algorithms.

In this work we are only using the magnitudes in the six (or four) available Rubin bands as inputs to the photometric estimation. We do not use other inputs, such as object size or morphology, or photometry from other surveys. As stated in Sec. 2.1.2, we systematically use the `gaap1p0` version of the fluxes, as these are considered the most reliable for color estimation.

3.1 Template fitting based estimators

RAIL’s template-based fitting algorithms (Lephare and BPZ, see references in Tab. 3 for more details) estimate photometric redshifts by comparing observed galaxy photometry to a set of predefined theoretical or empirical galaxy templates. These templates represent a range of galaxy spectral energy distributions (SEDs) that are meant to span the expected range of galaxy types observed in the Universe. Synthetic model fluxes are computed by red-shifting each SED to a set grid of values and convolving with the Rubin filter curves, which characterizes the expected variation in observed colors with redshift. Our algorithms calculate the chi-square/likelihood for each SED at each grid point by comparing the model fluxes in our photometric bands to the observed fluxes and uncertainties. This process enables the algorithm to determine the relative likelihood for each redshift for a given galaxy by identifying the template that best matches its observed color signature. An optional empirical Bayesian prior can also be applied to produce a posterior probability rather than a straight likelihood if information on the expected distributions are available. The training phase for these algorithms consist mostly of preparing pre-computed tables for the expected fluxes in each filter band for each SED template.

Two sets of templates are included in the `rail_base` software package and used in this note: the “CWWSB” templates that are the default templates used by `rail_bpz` and described in Coe et al. (2006), consisting of eight total SEDs, one Elliptical template, two Spiral templates, and five Irregular/Starburst template. In this note, these SEDs are used to compute synthetic colors that we compare to the observational data in Figures 3 and 4, but are otherwise not used. The second template set consists of those described in Ilbert et al. (2009), consisting of 31 synthetic SEDs (including some that are “interpolated” between adjacent templates by taking the mean of the two SEDs). These SEDs are included with `rail_lephare` as “COSMOS_mod.list”

in the `lephare-data` package. These SEDs are used by both template codes `rail_bpz` and `rail_lephare`. The 31 base SEDs contain no internal dust extinction, `rail_lephare` is configured to include dust automatically, while `rail_bpz` required dust to be added explicitly, and additional SED templates that include internal extinction added to the input SED list. LePhare also includes a set of stellar and AGN SEDs, which are available in the `lephare-data` package.

3.2 Machine learning based estimators

RAIL’s machine learning algorithms (CMNN, DNF, FlexZBoost, GPZ, KNN, TPZ) are described in detail in the publications listed in Tab. 3. In short these algorithms attempt to perform a regression analysis to estimate the redshift from the input photometric data. It is a well-known limitation of machine learning algorithms that they exhibit biases when presented with non-representative data, or are asked to extrapolate results outside region of their training data.

3.3 Analysis framework and bookkeeping software

The `rail_projects` software package within the RAIL ecosystem provides an essential framework for managing and organizing large-scale photometric redshift estimation workflows. It acts as a project management and book-keeping tool that helps users streamline their research, especially when working with complex or large datasets, like those associated with the Rubin DP1 data. The primary focus of `rail_projects` is to offer a systematic way to track different stages of data processing, model training, evaluation, and results across multiple experiments, ensuring that all tasks are well-documented and reproducible.

Additionally, `rail_projects` facilitates the management of large datasets by organizing data into structured directories and providing interfaces for batch processing. It allows users to scale up their work to handle not only large catalogs of objects but also multiple datasets and redshift estimation tasks. The package ensures that datasets and models are kept in sync across various stages, from raw input data to intermediate results to final outputs, and makes it easier to systematically export data coherent products.

3.4 Data exploration and algorithm optimization

We explored dozens of different configuration settings, covering analysis choices such as which flux measurements to use, the machine learning hyper-parameters, which sets of SED

templates to use, among others. The `dp1/dp1.yaml` configuration file (see Sec. A.1) captures the set of configurations that we tested and labels each set into an analysis “flavor” to facilitate bookkeeping.

After examining the performance of the algorithms under different configurations, we settled on a set of optimized parameters for each algorithm and on the `gaap_1p0` fluxes to produce the best-effort catalogs for DP1. These optimized configurations parameters are captured in the `dp1/dp1_v1.yaml` configuration file.

3.5 Production of redshift catalogs

We produced photo-*z* catalogs in two different software frameworks, both to test the software pipelines as well as to facilitate using associated tools to distribute the resulting data products.

3.5.1 Catalogs in the `rail_projects` framework

We used `rail_project` to run RAIL’s `EstimatePipeline`, to perform the photometric redshift estimation on the unlabeled data sets for all of the algorithms.

Specifically, we used `rail_project` to run all of the configuration flavors defined in the `dp1/dp1.yaml` and `dp1/dp1_v1.yaml` configuration files. The resulting data products are internally available as files at the USDF and NERSC as described in App. B.3. It is expected some of the products will be distributed via other methods in the near future (see App. B.2 and App. B.4).

3.5.2 Catalogs in the Rubin data management framework

We also used the `meas_pz` and `meas_pz_extensions` software packages to create photo-*z* catalogs in the Rubin data management framework.

Specifically, we imported trained models produced in the `rail_project` into the DP1 data Butler and produced photo-*z* estimates for every object in DP1 for each algorithm.

We did this for the `baseline` analysis flavor, (consisting of “out-of-the-box” algorithms, with all settings at default values) and the optimized 6 band (`dp1_optimize`) flavor defined in `dp1/dp1_v1.yaml`. The resulting data products are internally available via data Butlers at USDF and NERSC as de-

scribed in App. B.1. As with the `rail_projects` created catalogs, it is expected some of the products will be distributed via other methods in the near future (see App. B.2 and App. B.4).

4 Performance

We evaluated all the algorithms listed in Tab. 3 for both scientific and technical performance.

4.1 Redshift estimator performance

We primarily evaluated the performance using the “test_v1” catalog which consist of 2,437 galaxies randomly drawn from the cross-match between the DP1 object catalog in ECDFS with the reference redshifts described in Sec. 2.2. Additionally, we cross matched the DP1 object catalog in the **SV_38** field with the DESI DR1 BGS, LRG and ELG galaxies as a secondary validation set (“test_DESI”).

Specifically, we ran the PzPipeline (Inform → Estimate → Evaluate) to train models for per-object $p(z)$ estimation on the training data set, use those models to obtain photo- z estimates for the test data set, and then evaluate the performance of the photometric point-estimate of the redshift against the spectroscopic estimates.

For each algorithm, we produce a set of performance monitoring plots, described in Sec. A.6. In Fig 9, we show examples of these performance monitoring plots for two algorithms: `knn` and `bpz` for the 6-band models and using the mode for the specific point estimate. Fig. 10 shows the performance monitoring of `knn` and `bpz` when 4-bands are used. In the scattering plot, the mean, standard deviation, 3σ outliers and absolute outliers of

$$\Delta z = \frac{z_{\text{phot}} - z_{\text{spec}}}{1 + z_{\text{spec}}} \quad (1)$$

are shown in legend. We summarized these statistics in Table 4.

4.1.1 Redshift estimation quality flags

Applying a simple cut on the RMS of the $p(z)$ distribution can dramatically reduce the catastrophic outlier rate. In Fig. 11 we show how the efficiency and purity obtained on the test

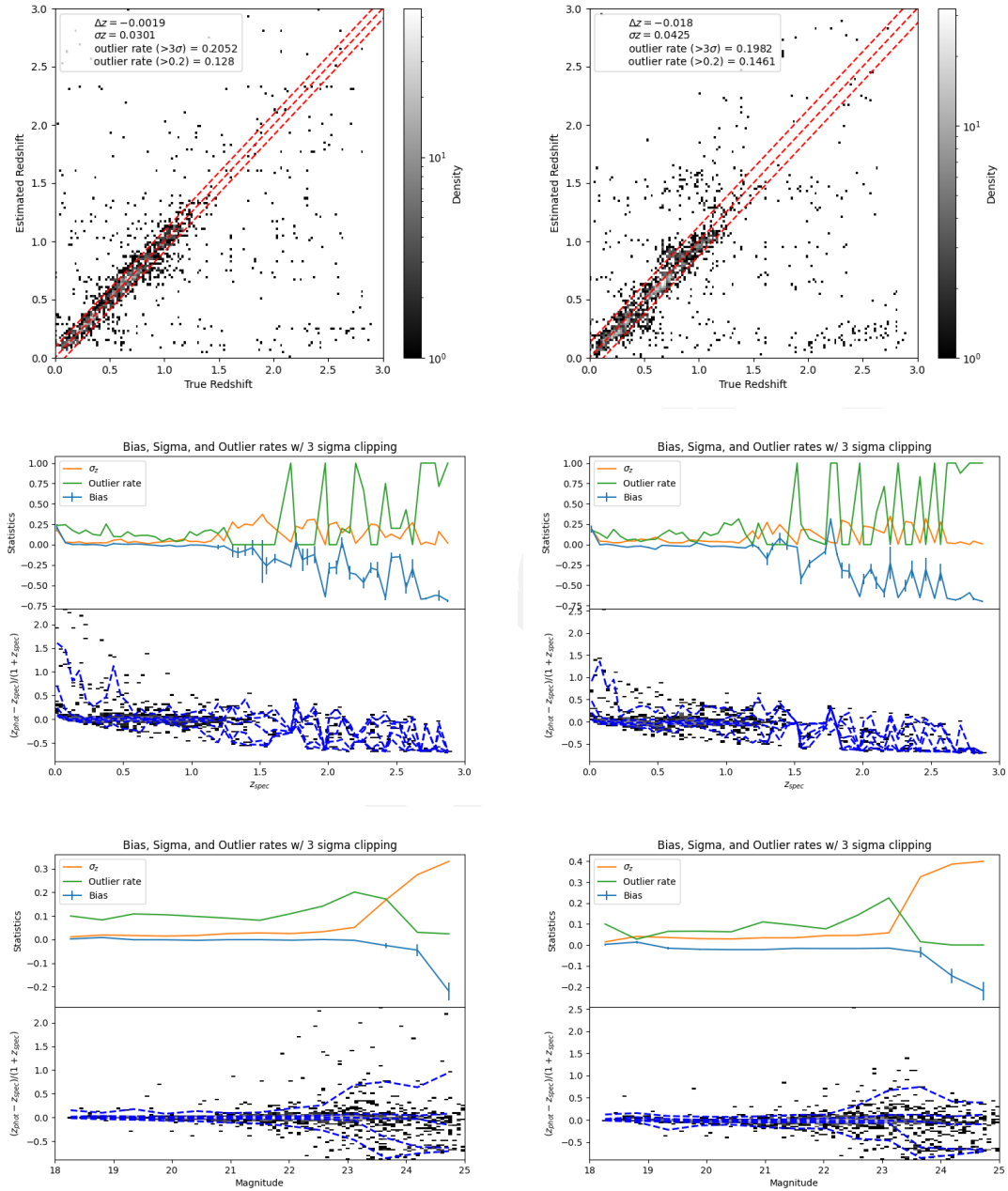


FIGURE 9: Estimator performance on the “gold” (six-band) dataset. Top row: photometric point estimate of the redshift versus spectroscopic redshift for KNN (left) and bpz (right). Middle row: performance metrics, (width and bias of residual, and outlier rate) versus spectroscopic redshift. Bottom row, performance metrics versus i-band magnitude.

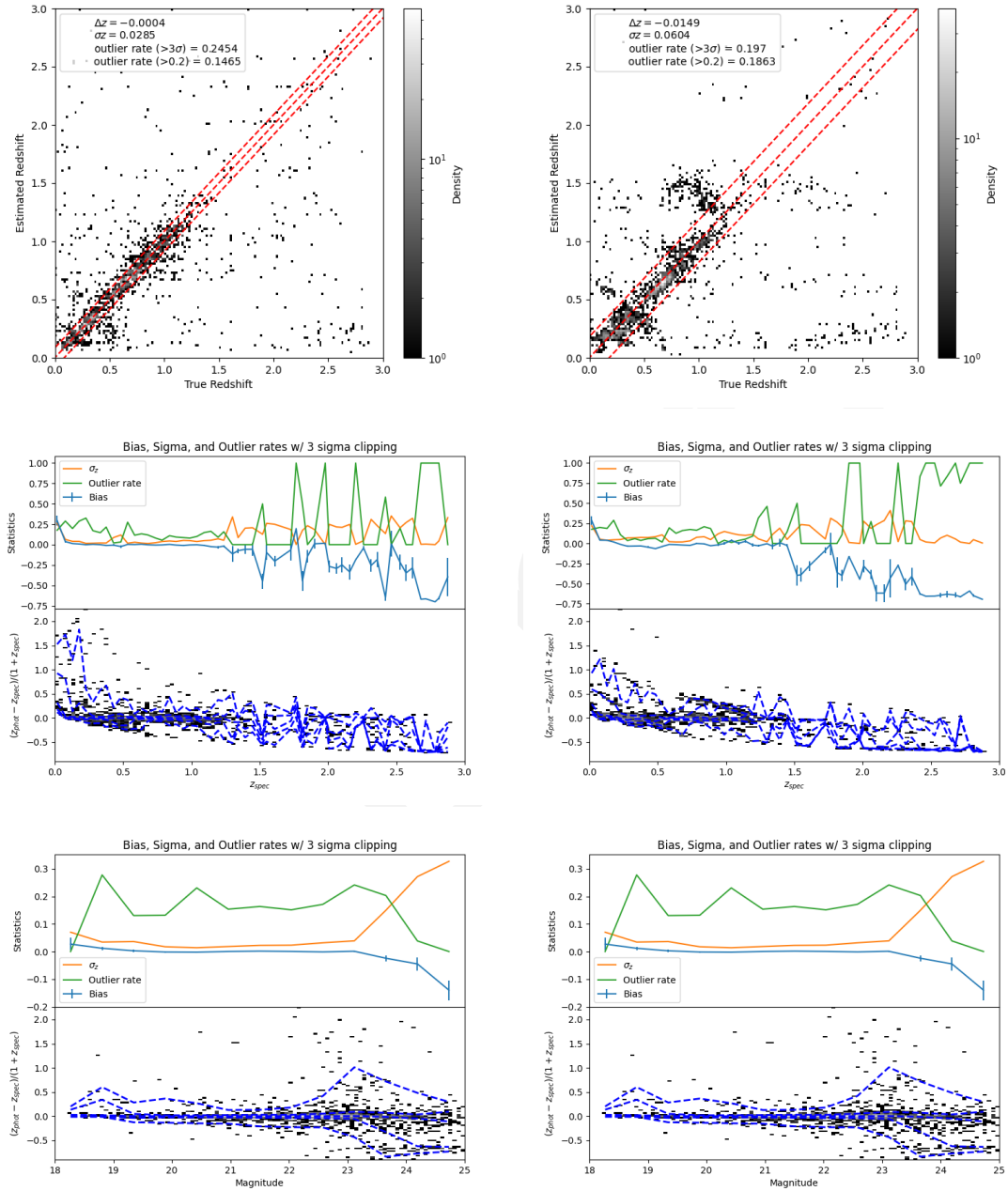


FIGURE 10: Estimator performance on the “gold_4_band” (four-band) dataset. Top row: photometric point estimate of the redshift versus spectroscopic redshift for KNN (left) and bpz (right). Middle row: performance metrics, (width and bias of residual, and outlier rate) versus spectroscopic redshift. Bottom row, performance metrics versus i-band magnitude.

TABLE 4: Performance metrics for photo- z algorithms using 6-band data, with 4-band results shown in parentheses.

Algorithm	Bias	σ	3σ Outlier Rate	$\Delta z > 0.2$ Outlier Rate
FlexZBoost	0.000 (0.000)	0.0246 (0.0274)	0.215 (0.221)	0.113 (0.124)
kNN	-0.002 (-0.000)	0.0301 (0.0285)	0.205 (0.245)	0.128(0.147)
CMNN	0.000 (-0.002)	0.0369 (0.0729)	0.227 (0.212)	0.160 (0.226)
DNF	-0.002 (-0.001)	0.041 (0.0327)	0.189 (0.215)	0.138 (0.121)
TPZ	-0.001 (-0.002)	0.050 (0.0524)	0.154 (0.156)	0.117 (0.127)
GPz	0.032 (0.018)	0.166 (0.106)	0.056 (0.110)	0.260 (0.198)
BPZ	-0.018 (-0.015)	0.0425 (0.060)	0.198 (0.197)	0.146 (0.186)
LePhare	-0.012 (-0.015)	0.0344 (0.0699)	0.207 (0.175)	0.139 (0.185)

sample varies with a single additional cut, x in $\sigma_{p(z)} < x$. In Fig. 12 we show how applying a cut at $\sigma_{p(z)} < 0.15$ improves the scatter of the photometric redshift point-estimate versus spectroscopic redshift distribution. As fainter objects and objects at higher redshift often have larger $p(z)$ RMS values, cuts on RMS will likely result in relative shifts in the magnitude and redshift distribution depending on the cut value.

4.1.2 Validation in the SV_38 field

We cross-matched our photo- z estimates with the DESI DR1 LSS galaxies (BGS, LRG, and ELG) in the **SV_38** field. This cross-matched catalog can serve as additional validation for our photo- z results. In Fig. 13, we show the median photometric redshift estimates versus the DESI spectroscopic redshifts all eight tested photo- z algorithms. Each panel corresponds to a different algorithm, allowing us to visually assess the bias, scatter, and potential outlier behavior across a wide redshift range. Overall, the agreement with DESI DR1 spectroscopic redshifts provides a robust sanity check on the performance of our methods in the early LSST data regime.

4.2 Technical performance

Tab. 5 shows compute times and file sizes for both the training and estimation phases of the analysis. The metrics from the training phase were taken from running the algorithm's "Inform" stages on USDF on the "training_v1" dataset of 7,000 objects. The metrics from the estimation phase were taken from algorithm's "Inform" stages on USDF on the "test_v1" datasets of 2,437 objects.

We note that the technical performance requirements depend heavily on the scientific use

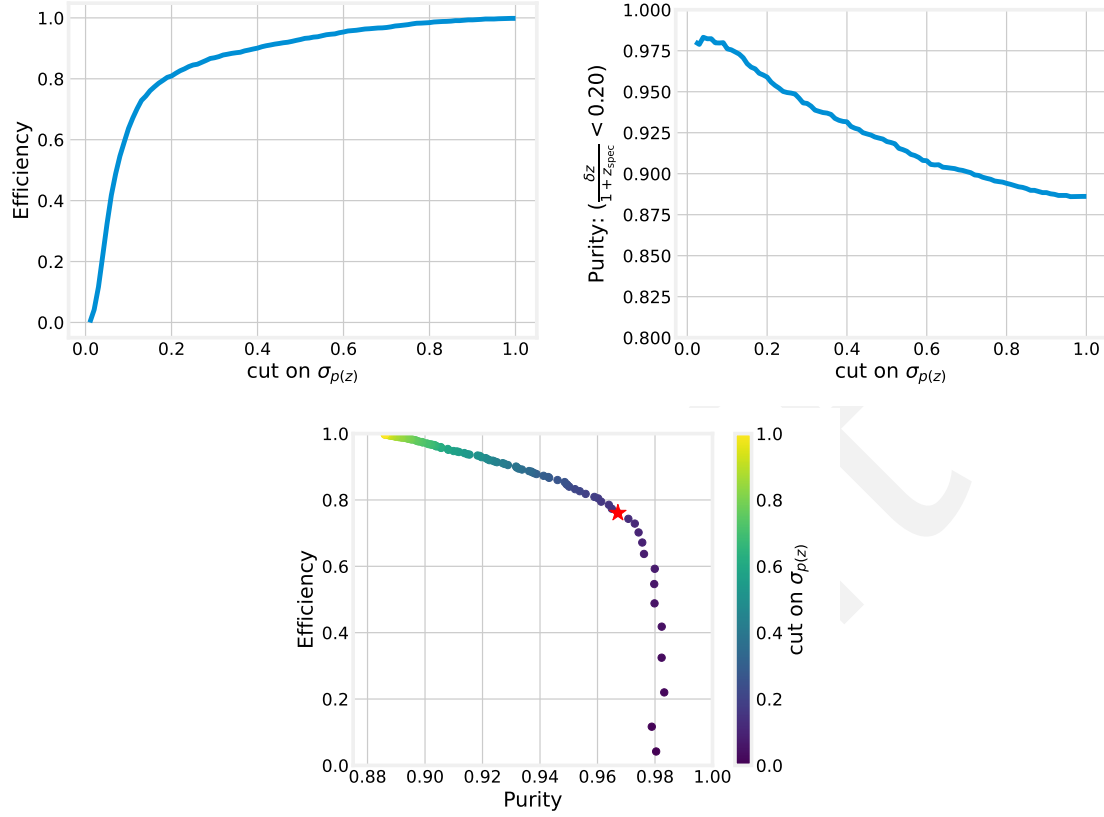


FIGURE 11: Efficiency (top left) and “purity”, i.e., fraction of objects with $\frac{\delta z}{1+z_{\text{spec}}} < 0.20$ (top right) versus quality cut, x in $\sigma_{p(z)} < x$. Bottom, purity version efficiency with cut value shown by the color scale. The red star shows the point for a cut value $\sigma_{p(z)} < 0.15$.

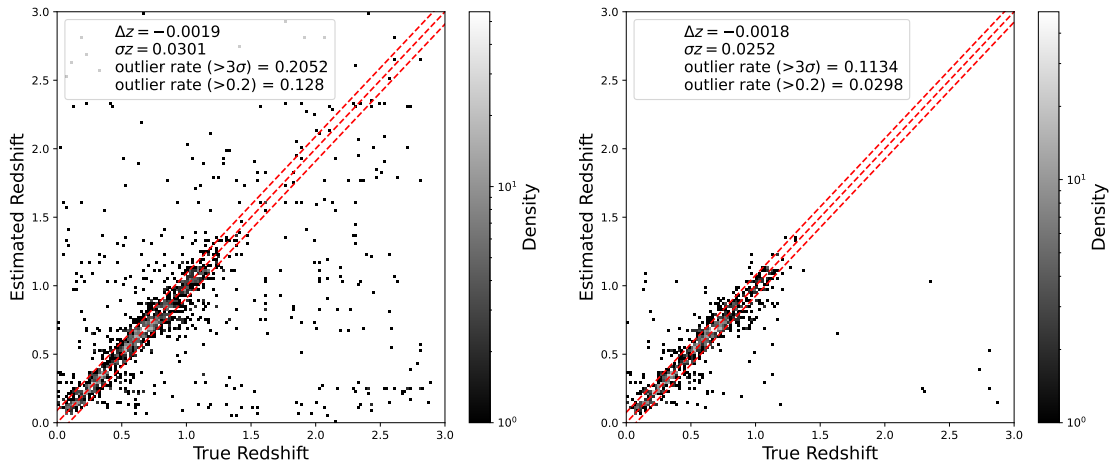


FIGURE 12: Estimator performance for TPZ without (left) and with (right) a quality cut of $\sigma_{p(z)} < 0.15$.

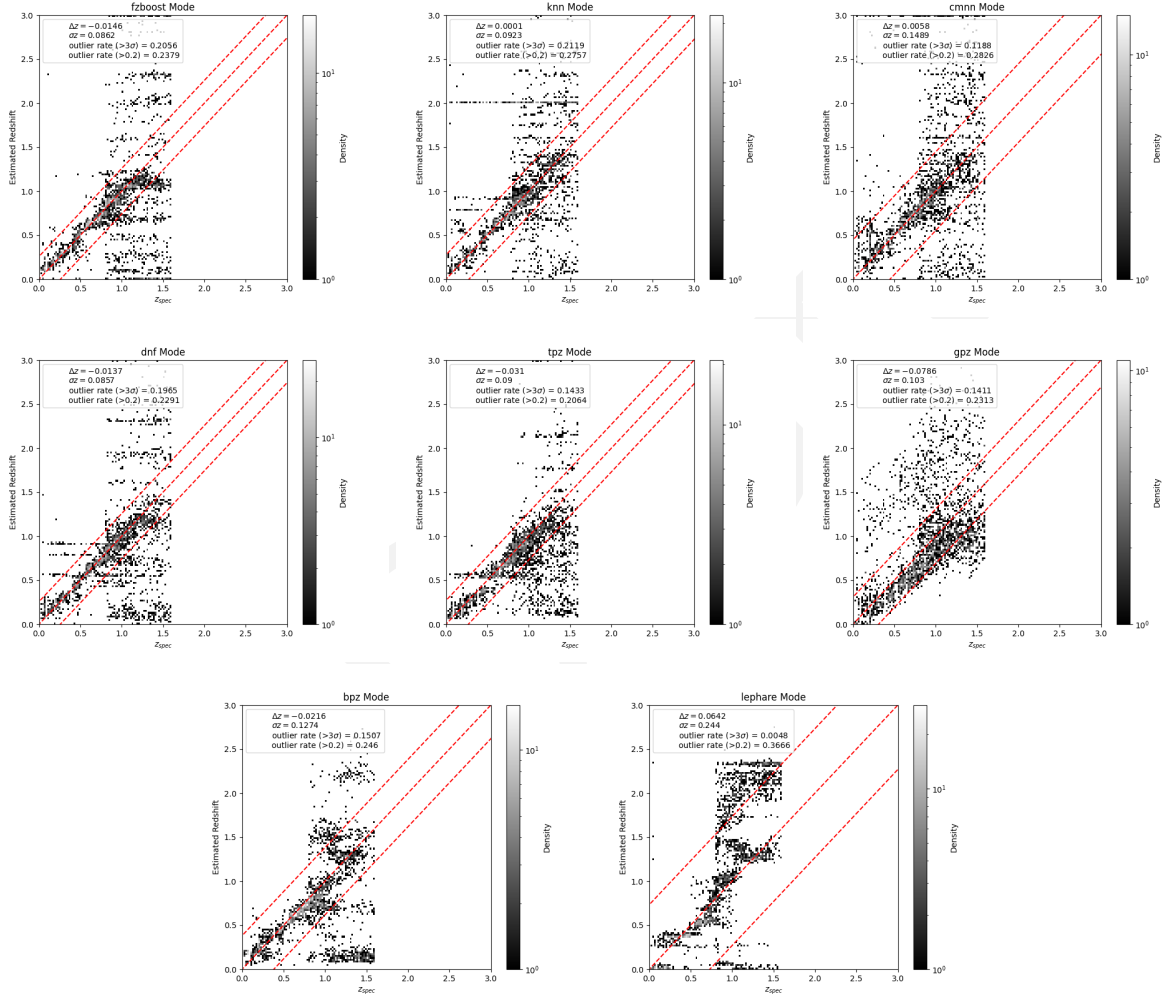


FIGURE 13: Mode photometric redshift estimates versus DESI DR1 spectroscopic redshifts for galaxies in the **SV_38** cross-matched sample. Each panel shows results from a different photo- z algorithm: FlexZBoost, kNN, CMNN, DNF, TPZ, GPz, BPZ, and LePhare (from left to right, top to bottom). These comparisons provide an external validation of photo- z performance using Large-Scale Structure spectroscopic redshifts samples DESI.

Algorithm name	Training		Estimation	
	Time [s]	Model Size [MB]	Speed [ms / object]	Data Size [b / object]
Project Supported				
BPZ	< 5	1	1.5-10	~ 2400
KNN	8 - 30	1	0.6-1.5	~ 240
Community Supported				
CMNN	< 5	1	2.4	54
DNF	< 5	2	0.6	~ 2400
FlexZBoost	55 - 100	35 - 50	4-13	~ 2400
GPz	10 - 15	1	< 0.5	32
LePHARE	300 - 600	1	50 - 180	~ 2400
TPZ	5-20	7 - 300	7 - 120	~ 2400

TABLE 5: Algorithm compute times and files sizes. The algorithms were trained on 7000 objects and evaluated on 2,437 objects. The variations shown reflect the differences in processing times with different sets of hyper-parameters.

case. For example, in supernova cosmology, the number of objects will be in the 100's of thousands, while for weak-lensing cosmology, it will be in the billions. Accordingly, these give very different constraints on the required processing speed of the photo- z estimation. While seconds per object is supportable for supernova cosmology, weak-lensing will require processing times in the milliseconds (ms) per object range.

5 Limitations and Caveats

The photo- z estimates described in this note were done on a best-effort basis by members of the "Photo- z Science Unit" and are not official Rubin data products. As such, the release of the various DP1-related data products comes with a number of caveats.

- The size and depth of the training set seriously limits the robustness of the training past a redshift of $z \simeq 1.5$. See in particular Fig. 2.
- In general, photo- z algorithms do not perform well with objects with marginal detections or with non-detections in several bands. Simply put, the photometric uncertainties in such objects often overwhelms the information that is present and the photometric estimates are very uncertain. See Fig. 14 for an example of the PDF of such an object. Similarly, see, Fig 15 as an illustration of how the accuracy of a photo- z point-estimate degrades significantly for faint objects.

- Taking points (1) and (2) together, we do not advise to use any of the provided DP1 photo- z estimates without applying cuts on the detection significance or the width of the $p(z)$ distribution, or both.
- While we did put some work into optimizing the performance of the various algorithms, this was by no means comprehensive, and we urge caution in drawing any conclusions about the relative merits of the algorithms.
- The training and test sets were drawn from the same set of spectroscopically matched objects. As such, the training set is fairly representative of the test set. On the other hand, the full DP1 catalog does not have any spectroscopic selections applied, so the training and test sets will not be representative of the full DP1 data set.

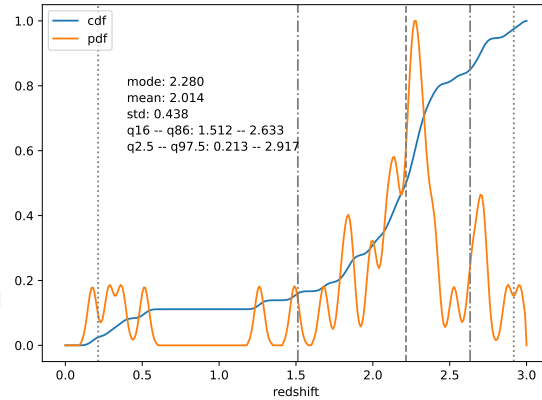


FIGURE 14: Example of a $p(z)$ distribution for a faint object is not detected in all bands. (In this case the object was only significantly detected in 'g' and 'r' bands).

6 Summary and Conclusion

This note describes an initial calculation of photometric redshifts on the Rubin Data Preview 1 dataset using the tools and workflows developed specifically for this purpose by the Photo- z Science Unit. These early results are very promising, we show that we can successfully match Rubin data with spectroscopic samples in the ECDFS field with deep six-band Rubin coverage to create a reference sample and run our software pipelines to generate a catalog of individual object redshifts consisting of both point 1D redshift PDFs and point estimate redshifts. We show results for eight photo- z algorithms using a reserve set of redshifts, and see that photo- z performance is in-line with expectations both in terms of our redshift predictions and for compute times. We also test our estimates against an independent set of redshifts from the

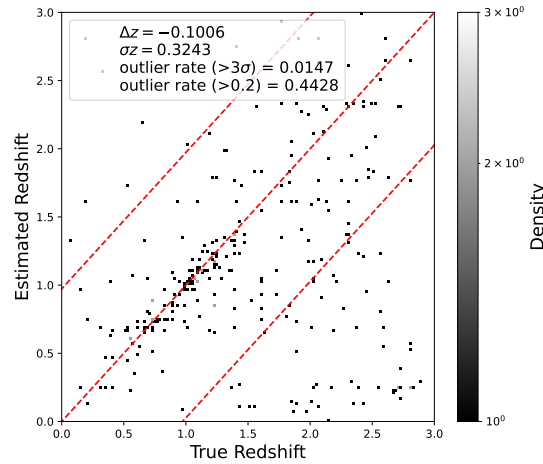


FIGURE 15: Point-estimates versus spectroscopic redshifts for all galaxies in the reserved test set with $m_i > 23.5$, showing the degradation in estimation performance for faint objects.

SV_38 field with shallower coverage and imaging in only four of the six Rubin bands. We describe available data products and anticipated methods to access them. Overall, this is a very successful initial test of photo- z pipelines.

However, work will continue on several fronts to further optimize the performance of our redshift predictions. As mentioned earlier in the note, the definition of the true redshift reference sample used to train our algorithms has a large impact on results, being the “truth” used to define the flux/color-to-redshift mapping that is inherent to photo- z estimation. We will work to refine our reference sample definition, which objects to include/exclude based on quality flags, explore the tradeoffs of including grism and many-band photo- z estimates with larger redshift uncertainties in our training sample, and other such effects. As more and more Rubin data is taken, areal coverage is increasing, including coverage of additional deep calibration fields with existing spec- z datasets. This will naturally improve photo- z performance. Expansion of reference redshifts will also enable the development of improved object flagging for identifying which redshifts are trustworthy and those which are likely to be incorrect.

We will also continue to examine the Rubin photometry and evaluate the performance of multiple measurement algorithms. In this note we used Gaussian Aperture (Gaap) fluxes and magnitudes, as they are expected to produce consistent colors. We will undertake a thorough exploration of multiple flux measurements (e. g. cModel, Sersic, additional Gaap apertures) to test which combinations lead to the best photo- z estimates, as well as combinations like using multiple aperture magnitudes that may contain information on the galaxy light profile

that could help to constrain redshift. As we are still in the engineering and testing phase of the project, there is ongoing work to characterize system performance, and improved understanding of the system will likely lead to better photo- z performance. For example, there may be some hints that u-band fluxes are overestimated and u-band uncertainties underestimated at faint magnitudes (see the high redshift u-g colors in Fig. 3 and figures in Vera C. Rubin Observatory (RTN-095)). DP1 data was obtained using LSSTComCam, while future data releases will use LSSTCam, and thus performance for DP2 and beyond may be shaped by slight differences between the two instruments.

In this analysis, while there was some optimization of code-specific parameters for the individual estimators, e. g. the number of neighbors used for KNN, the number of trees used for TPZ, it was not comprehensive. In addition, the optimal values will likely change slightly as we refine both our spectroscopic reference sample and our photometric inputs. A thorough exploration of the code parameters will happen as we converge on our final photometric and spectroscopic setup.

As stated in the introduction, we expect feedback from science users as they explore the data and discover issues that we have not anticipated, and that we will incorporate that feedback in future work.

Key takeaway: while very promising, these results are far from final, there is ongoing work to optimize photo- z performance. We have plans in place and will continue to refine as new data arrives, and these plans should lead to improved redshift performance.

A Data products

In the course of this work, we generated several data products, including redshift estimates and a variety of others that can be used to reproduce those estimates and to facilitate thorough evaluation of the algorithm performance. These products include: configuration files (Sec. A.1), ancillary inputs (Sec. A.2), trained models (Sec. A.3), redshift estimates (Sec. A.4), summary statistics (Sec. A.5, and performance monitoring plots (Sec. A.6), all of which are essential for understanding the quality of the photometric redshift estimates and for refining the algorithms. Together, these data products provide a comprehensive framework for conducting, managing, and evaluating photometric redshift estimation workflows in Rubin DP1. They ensure that the process is not only scientifically rigorous but also organized and reproducible, enabling effective collaboration and ongoing refinement of photometric redshift techniques.

A.1 Configuration files

The configuration files, collected in the GitHub `rail_project_config` repository, provide the necessary parameters and settings to control the various stages of the redshift estimation process, are a key element of `rail_projects` workflow. These files typically include specifications for the photometric bands used (e.g., `g`, `r`, `i`, `z`), the algorithm choices (e.g., template fitting or machine learning methods), details about data pre-processing, such as feature normalization and handling of missing data. These configuration files also specify the training and validation dataset splits, hyper-parameters for machine learning models, and paths for input/output data. These files are essential for ensuring reproducibility and for sharing the exact settings used in different redshift estimation runs, enabling other researchers to replicate or extend the analysis.

A.2 Ancillary input files

In addition to the configuration files, we require ancillary inputs such as the galaxy spectral energy distribution (SED) templates and the filter throughputs. Galaxy SED templates are collections of theoretical or observed spectra for galaxies at different redshifts and with different properties (e.g., galaxy type, age, star formation history). These templates are used by template-fitting algorithms to model the expected galaxy colors as a function of redshift, allowing for the estimation of photometric redshifts by comparing observed colors to those predicted by the templates. Two SED sets are employed in this note, they are each described

in Section 3.1.

The filter throughputs specify the characteristics of the fiducial observational filter transmission curves, and are used in Rubin DP1, including their corresponding central wavelengths. These filter curves are employed in calculating the synthetic fluxes or magnitudes expected from each of the SED templates used in template-fitting algorithms, and for matching observational data to predicted theoretical values, e. g. the predicted colors shown in Figure 3.

A.3 Estimator data models

After training, the trained models for each photometric redshift estimation algorithm are stored as serialized files, either in Pickle (for Python-based models) or YAML (for model configurations) format. These models encapsulate the learned relationships between photometric features (such as magnitudes and colors) and redshift values, allowing them to be applied to new data for redshift estimation. These files store the final state of the model, including the weights, biases, and other learned parameters for machine learning models. In the case of template-fitting methods, the corresponding model files may include the template sets and the fitting parameters. The Pickle or YAML format ensures that the models can be easily loaded, applied to new datasets, and evaluated in future studies.

A.4 Redshift estimates stored as QP ensembles

The per-object redshift estimates generated by the photometric redshift algorithms are stored in qp files. These files serve as containers for storing the redshift predictions for each object in the dataset before any detailed statistical analysis or final reporting. In the qp files, each galaxy's redshift estimate is stored in a format that is compatible with the algorithm providing the estimate. Specifically, for each object we store distribution $p(z)$, which, depending on the algorithm, maybe represent a posterior probability, a likelihood, a conditional likelihood, or just a hunch as to redshift of the object in question. These qp files are designed to be lightweight and easy to query, allowing users to quickly retrieve the redshift estimates for individual objects. The structure of qp files is optimized for efficient access and data retrieval, making it easier for users to process and analyze large numbers of photometric redshift estimates across large datasets like Rubin DP1. Additional information, including usage examples, about qp and qp files is available at <https://qp.readthedocs.io/en/main/>.

A.5 Per-Object Point Estimates

In addition to the raw redshift estimates, the qp files also store per-object point estimates in the form of an ancillary table. These estimates include crucial information about the quality and reliability of each photometric redshift estimate, such as the uncertainty in the redshift prediction (e.g., the confidence interval or standard deviation or, the likelihood score (in the case of probabilistic models). This table will eventually includes flags for identifying objects with low-confidence estimates or those that may be outliers. These summary statistics are important for evaluating the overall performance of the redshift estimation algorithms on an object-by-object basis and are often used to filter out low-quality or problematic redshift estimates before conducting larger statistical analyses.

We also generate meta catalog file that collects the Object ID, magnitude information, and point estimates. The point estimates are named by “{algorithm}_z_{point estimate name}”. Here is a list of the point estimates and summary statistics that are stored for each algorithm and the corresponding column names:

1. **mean**: is the per-object expectation of redshift given by the PDF.
2. **median**: is the 50% percentile of the PDF.
3. **mode**: is the redshift that yields maximum PDF evaluated on 301 grid points between $z = 0-3$.
4. **err68_lower(upper)**: is the 16th (84th) percentile of the per-galaxy PDF, which correspond to the 1σ confidence interval.
5. **err95_lower(upper)**: is the 2.5th (97.5th) percentile of the per-galaxy PDF, which correspond to the 2σ confidence interval.

Fig. 16 shows an example of the $p(z)$ distribution, point-estimates and summary statistics for a single object.

Yes.

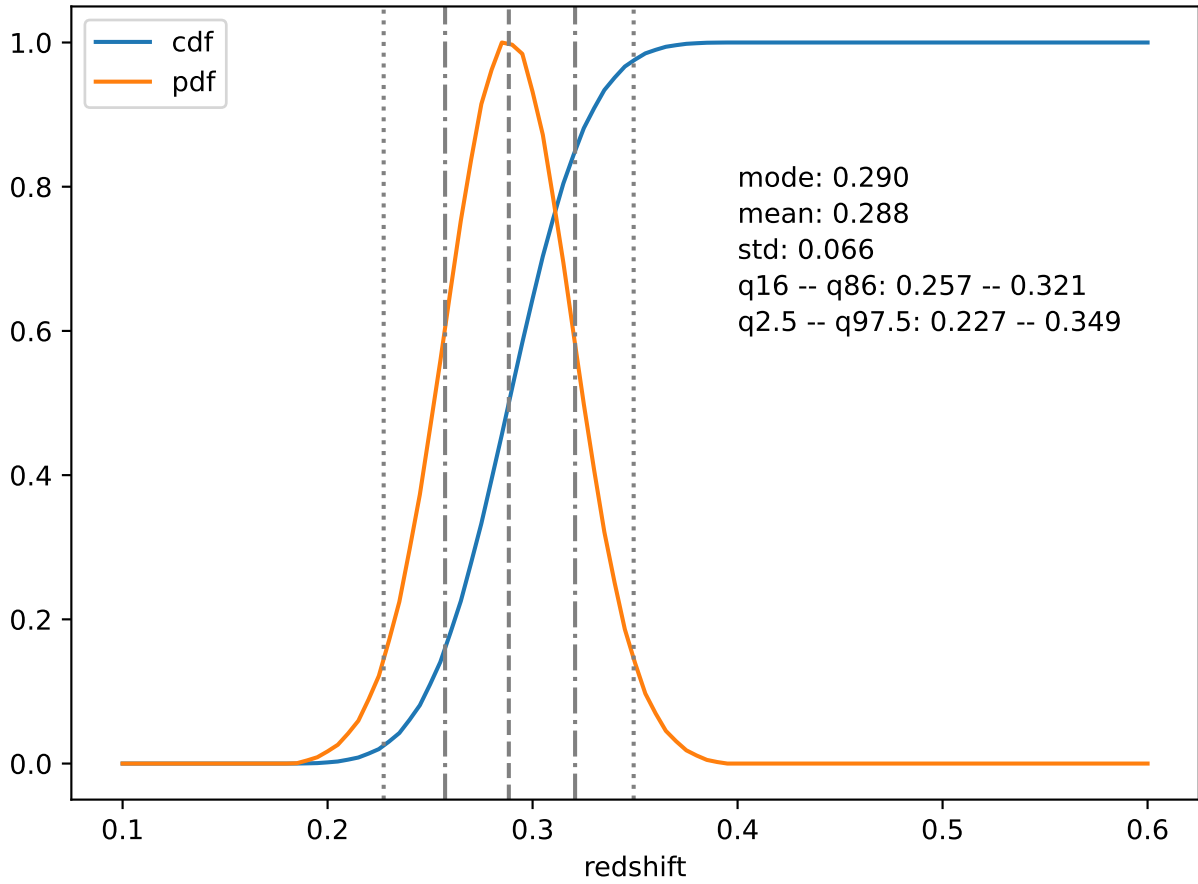


FIGURE 16: Single object $p(z)$ estimate, showing both the PDF, and the CDF, as well as the summary statistics described in the text.

A.6 Performance Monitoring Plots

Finally, we produced standardized performance monitoring plots as part of the photometric redshift workflow. These plots provide visual representations of the model's performance, allowing users to assess how well the redshift estimates align with the true redshifts from the spectroscopic training data. Common plots include:

- Input dataset characterization plots, as shown in Sec. 2.1.2.
- Scatter plots to visualize the relationship between the predicted and true redshifts, highlighting any systematic biases or non-linearities.
- Redshift comparison plots, e.g., the bias or width of the redshift residuals, $\frac{z_{\text{phot}} - z_{\text{spec}}}{1 + z_{\text{spec}}}$, as a function of redshift or object magnitude.

To robustly summarize the distribution of residuals while minimizing the influence of outliers, we compute biweighted statistics following a two-step procedure. First, we apply an iterative 3σ clipping to the input array using `scipy.stats.sigmaclip`, repeating the clipping process `nclip` times (default is 3). This removes extreme outliers before computing robust estimators. We then calculate the *biweight location* and *biweight scale* of the clipped subset using the `astropy.stats` functions, which yield robust analogs to the mean and standard deviation, respectively.

In addition, we compute two outlier rates: the *relative outlier rate*, defined as the fraction of values in the original sample that deviate from zero by more than $3\times$ the biweight scale, and the *absolute outlier rate*, defined as the fraction of values exceeding a fixed threshold set by `self.config.abs_out_thresh`. This framework provides both robust central moments and a diagnostic of extreme deviations in the full sample.

Examples of the two latter types of plots for the BPZ and KNN algorithms are shown in Sec. 4

B Data Distribution

To support both the Rubin community and the DESC, we will distribute these data products in several different ways:

1. Via the Rubin Data Butler (see Sec. B.1).
2. Via the Photo-z Server (see Sec. B.2).
3. Directly as files at the USDF and NERSC (see Sec. B.3).
4. Via LSDB (see Sec. B.4)

In each of these distribution mechanisms there are a number of metadata that are used as fields in defining specific data products. In particular:

- **algo**: specifies a particular photo-z estimation algorithm (see Tab. 3),
- **selection**: specifies a particular data selection (see Sec. 2.1.1),
- **flavor**: specifies a particular set of configuration parameters,
- **dataset**: specifies a particular dataset (see Sec. 1).
- **version**: specifies a processing version (e.g., v1, v2).

B.1 Distribution via the Rubin Data Butler

Creation of photometric redshift estimates using RAIL in the Rubin DM framework and distribution via the Rubin Data Butler is supported the `meas_pz` software package for DM supported algorithms and by the `meas_pz_extensions` software package for the community-supported algorithms described in this note.

For DP1, the following objects stored in the data butler at USDF and NERSC are listed in Tab. 6.

B.2 Distribution via the Photo-z Server

The Photo-z Server (or PZ Server) is a web-based service available for the LSST community to create and host PZ-related lightweight data products. It relies on the infrastructure of the Brazilian Independent Data Access Center and is developed and maintained by LIneA as part of the Brazilian in-kind contribution program.

Dataset type	Description
Collection: pretrained_models/pz/DP1/{selection}/{flavor}	
pzModel_{algo}	Estimator models (see A.3)
Collection: LSSTComCam/runs/DRP/DP1/pz/DM-51523/{selection}/{flavor}/{version}	
pz_estimate_{algo}	QP ensembles (see A.4)
pz_{algo}_config	Configuration parameters (see A.1)
pz_{algo}_log	Log files
pz_{algo}_metadata	Processing metadata

TABLE 6: Photo-z related objects stored in the Rubin Data Butler. Note that {selection} describes the selection applied to the trained data set.

The PZ server is open to any LSST member with a valid RSP account without the need for an extra local registry. Users can log into the system simply using the RSP credentials.

All data products described in this document will be hosted on the PZ Server, along with their respective metadata and documentation. There are two ways to access these data products:

B.2.1 From the PZ Server website

Data products are listed both on the 'Rubin PZ Data Products' page (for official data products released or recommended by LSST DM) and 'User-generated data products' for data products produced or uploaded by the LSST community members. The DP1 PZ data products described in this document will be available in the first one.

B.2.2 Via the pzserver Python library

Similar to the RSP, the PZ Server provides an API interface that enables users to access data through Python scripts from any location, provided they know the product name as registered on the server.

If it is the first time using the library, it must be installed via pip in the terminal or in a notebook cell: `pip install pzserver`

Then, the `PzServer` class opens the remote connection to the PZ Server database. An access token is required for authentication. The token can be generated by users on the PZ Server website (top right corner menu on the home page).

```
from pzserver import PzServer
pz_server = PzServer(token="<paste your access token here>")
```

To display the product metadata and download it to the local working directory (if not in a Jupyter notebook, replace `display` for `get`):

```
pz_server.display_product_metadata(<product_id>)
pz_server.download_product(product_id, save_in=".")
```

Alternatively, it is possible to load a table directly into memory as a Pandas DataFrame or Astropy Table. For instance, to load a training set:

```
training_set = pz_server.get_product(<training_set_id>)
training_set.display_metadata()
```

A tutorial notebook with examples for all `pzserver` methods is available on the `pzserver` library's repository on GitHub.

B.3 Distribution as files at USDF and NERSC

All of the test and training files, the outputs from runs used to optimize the model hyperparameters, as well as from the optimized models, the photo- z estimates for the entire DP1 dataset, are all available in `rail_projects`-managed shared project area at both NERSC and USDF. These data products are listed in Tab. 7.

B.4 Distribution via LSDB

The Large Survey DataBase (LSDB) will host the DP1 data on the USDF and the Canadian IDAC RSP. The DP1 data will be turned into Hierarchical Adaptive Tiling Scheme (HATS) format. LSDB enables researchers to load large datasets with limited memory, and fast cross match to other surveys like DESI DR1 and GAIA DR3.

Object type	Relative Path
Training data sets	data/train/*.hdf5
Test data sets	data/test/*.hdf5
In : pz/projects/dp1/pipelines	
Pipeline configurations	{pipeline}_{flavor}.yaml
In : pz/projects/dp1/data	
Estimator models	{selection}_{flavor}/model_inform_{algo}.pkl
QP ensembles (for test data)	{selection}_{flavor}/output_estimate_{algo}.hdf5
QP ensembles (for other data)	{selection}_{flavor}/{dataset}/output_estimate_{algo}.hdf5

TABLE 7: Photo-z related files in the rail_projects-managed shared project areas.

The photo-z point estimates for DP1 in the **ECDFS,EDFS, Rubin_SV_95_-25**, and **Rubin_SV_38_7** will be served as a single tabular catalog in the LSDB. The location on the USDF for this catalog is /sdf/data/rubin/shared/lsdb_commissioning/dp1_pz_hats.

C References

- Almosallam, I.A., Jarvis, M.J., Roberts, S.J., 2016, MNRAS, 462, 726 (arXiv:1604.03593), doi:10.1093/mnras/stw1618, ADS Link
- Arnouts, S., Cristiani, S., Moscardini, L., et al., 1999, MNRAS, 310, 540 (arXiv:astro-ph/9902290), doi:10.1046/j.1365-8711.1999.02978.x, ADS Link
- Balestra, I., Mainieri, V., Popesso, P., et al., 2010, A&A, 512, A12 (arXiv:1001.1115), doi:10.1051/0004-6361/200913626, ADS Link
- Benítez, N., 2000, ApJ, 536, 571 (arXiv:astro-ph/9811189), doi:10.1086/308947, ADS Link
- Blake, C., Amon, A., Childress, M., et al., 2016, MNRAS, 462, 4240 (arXiv:1608.02668), doi:10.1093/mnras/stw1990, ADS Link
- Carrasco Kind, M., Brunner, R.J., 2013, MNRAS, 432, 1483 (arXiv:1303.7269), doi:10.1093/mnras/stt574, ADS Link
- Coe, D., Benítez, N., Sánchez, S.F., et al., 2006, AJ, 132, 926 (arXiv:astro-ph/0605262), doi:10.1086/505530, ADS Link
- Collaboration, D., Abdul-Karim, M., Adame, A.G., et al., 2025, Data release 1 of the dark energy spectroscopic instrument, URL <https://arxiv.org/abs/2503.14745> (arXiv:2503.14745)
- Colless, M., Dalton, G., Maddox, S., et al., 2001, MNRAS, 328, 1039 (arXiv:astro-ph/0106498), doi:10.1046/j.1365-8711.2001.04902.x, ADS Link
- Cool, R.J., Moustakas, J., Blanton, M.R., et al., 2013, ApJ, 767, 118 (arXiv:1303.2672), doi:10.1088/0004-637X/767/2/118, ADS Link
- De Vicente, J., Sánchez, E., Sevilla-Noarbe, I., 2016, MNRAS, 459, 3078 (arXiv:1511.07623), doi:10.1093/mnras/stw857, ADS Link
- D'Eugenio, F., Cameron, A.J., Scholtz, J., et al., 2025, ApJS, 277, 4 (arXiv:2404.06531), doi:10.3847/1538-4365/ada148, ADS Link
- Garilli, B., McLure, R., Pentericci, L., et al., 2021, A&A, 647, A150 (arXiv:2101.07645), doi:10.1051/0004-6361/202040059, ADS Link

- Graham, M.L., Connolly, A.J., Ivezić, Ž., et al., 2018, *AJ*, 155, 1 (arXiv:1706.09507), doi:10.3847/1538-3881/aa99d4, ADS Link
- [DMTN-049]**, Graham, M.L., Bosch, J., Guy, L.P., The DM System Science Team., 2022, *A Roadmap to Photometric Redshifts for the LSST Object Catalog*, Data Management Technical Note DMTN-049, Vera C. Rubin Observatory, URL <https://dmtn-049.lsst.io/>
- Hahn, C., Wilson, M.J., Ruiz-Macias, O., et al., 2023, *AJ*, 165, 253 (arXiv:2208.08512), doi:10.3847/1538-3881/acff8, ADS Link
- Helou, G., Madore, B.F., Schmitz, M., et al., 1991, In: Albrecht, M.A., Egret, D. (eds.) *Databases and On-line Data in Astronomy*, vol. 171 of *Astrophysics and Space Science Library*, 89–106, doi:10.1007/978-94-011-3250-3_10, ADS Link
- Huchra, J.P., Macri, L.M., Masters, K.L., et al., 2012, *ApJS*, 199, 26 (arXiv:1108.0669), doi:10.1088/0067-0049/199/2/26, ADS Link
- Ilbert, O., Capak, P., Salvato, M., et al., 2009, *ApJ*, 690, 1236 (arXiv:0809.2101), doi:10.1088/0004-637X/690/2/1236, ADS Link
- Izbicki, R., Lee, A.B., 2017, *Electronic Journal of Statistics*, 11, 2800, URL <https://doi.org/10.1214/17-EJS1302>, doi:10.1214/17-EJS1302
- Jenness, T., Bosch, J.F., Salnikov, A., et al., 2022, In: *Software and Cyberinfrastructure for Astronomy VII*, vol. 12189 of *Society of Photo-Optical Instrumentation Engineers (SPIE) Conference Series*, 1218911 (arXiv:2206.14941), doi:10.1117/12.2629569, ADS Link
- Jones, D.H., Read, M.A., Saunders, W., et al., 2009, *MNRAS*, 399, 683 (arXiv:0903.5451), doi:10.1111/j.1365-2966.2009.15338.x, ADS Link
- Kodra, D., Andrews, B.H., Newman, J.A., et al., 2023, *ApJ*, 942, 36 (arXiv:2210.01140), doi:10.3847/1538-4357/ac9f12, ADS Link
- Kriek, M., Shapley, A.E., Reddy, N.A., et al., 2015, *ApJS*, 218, 15 (arXiv:1412.1835), doi:10.1088/0067-0049/218/2/15, ADS Link
- Le Fèvre, O., Vettolani, G., Garilli, B., et al., 2005, *A&A*, 439, 845 (arXiv:astro-ph/0409133), doi:10.1051/0004-6361:20041960, ADS Link
- Lidman, C., Tucker, B.E., Davis, T.M., et al., 2020, *MNRAS*, 496, 19 (arXiv:2006.00449), doi:10.1093/mnras/staa1341, ADS Link

- Merlin, E., Castellano, M., Santini, P., et al., 2021, A&A, 649, A22 (arXiv:2103.09246), doi:10.1051/0004-6361/202140310, ADS Link
- Merlin, E., Santini, P., Paris, D., et al., 2024, A&A, 691, A240 (arXiv:2409.00169), doi:10.1051/0004-6361/202451409, ADS Link
- Momcheva, I.G., Brammer, G.B., van Dokkum, P.G., et al., 2016, ApJS, 225, 27 (arXiv:1510.02106), doi:10.3847/0067-0049/225/2/27, ADS Link
- NSF-DOE Vera C. Rubin Observatory, 2025, Legacy Survey of Space and Time Data Preview 1: Object searchable catalog, URL <https://www.osti.gov/servlets/purl/2570325>
- Quintana, H., Ramirez, A., 1995, ApJS, 96, 343, doi:10.1086/192122, ADS Link
- Raichoor, A., Moustakas, J., Newman, J.A., et al., 2023, AJ, 165, 126 (arXiv:2208.08513), doi:10.3847/1538-3881/acb213, ADS Link
- [PSTN-019]**, Rubin Observatory Science Pipelines Developers, 2025, *The LSST Science Pipelines Software: Optical Survey Pipeline Reduction and Analysis Environment*, Project Science Technical Note PSTN-019, Vera C. Rubin Observatory, URL <https://pstn-019.lsst.io/>
- Schlafly, E.F., Finkbeiner, D.P., 2011, ApJ, 737, 103 (arXiv:1012.4804), doi:10.1088/0004-637X/737/2/103, ADS Link
- Tasca, L.A.M., Le Fèvre, O., Ribeiro, B., et al., 2017, A&A, 600, A110 (arXiv:1602.01842), doi:10.1051/0004-6361/201527963, ADS Link
- The RAIL Team, van den Busch, J.L., Charles, E., et al., 2025, arXiv e-prints, arXiv:2505.02928 (arXiv:2505.02928), doi:10.48550/arXiv.2505.02928, ADS Link
- [RTN-095]**, Vera C. Rubin Observatory, 2025, *The Vera C. Rubin Observatory Data Preview 1*, Technical Note RTN-095, Vera C. Rubin Observatory, URL <https://rtn-095.lsst.io/>
- Zhou, R., Dey, B., Newman, J.A., et al., 2023, AJ, 165, 58 (arXiv:2208.08515), doi:10.3847/1538-3881/aca5fb, ADS Link

D Acronyms

Acronym	Description
---------	-------------

1D	One-dimensional
3D	Three-dimensional
AGN	Active Galactic Nuclei
API	Application Programming Interface
CDF	Cumulative Distribution Function
COSMOS	Cosmic Evolution Survey
DESC	Dark Energy Science Collaboration
DESI	Dark Energy Spectroscopic Instrument
DM	Data Management
DMTN	DM Technical Note
DP1	Data Preview 1
DP2	Data Preview 2
DR1	Data Release 1
DR3	Data Release 3
DRP	Data Release Processing
ECDFS	Extended Chandra Deep Field-South Survey
EDFS	Euclid Deep Field South
ELG	Emission-Line Galaxies
GOODS	The Great Observatories Origins Deep Survey
HST	Hubble Space Telescope
IDAC	Independent Data Access Center
JWST	James Webb Space Telescope (formerly known as NGST)
LRG	Luminous Red Galaxies
LSS	Large Scale Structure
LSST	Legacy Survey of Space and Time (formerly Large Synoptic Survey Telescope)
LSSTCam	LSST Science Camera
LSSTComCam	Rubin Commissioning Camera
MB	MegaByte
NED	NASA/IPAC Extragalactic Database
NERSC	National Energy Research Scientific Computing Center
PDF	Portable Document Format
PSTN	Project Science Technical Note
PZ	photo-z

RMS	Root-Mean-Square
RSP	Rubin Science Platform
RTN	Rubin Technical Note
SE	System Engineering
SED	Spectral Energy Distribution
SNR	Signal to Noise Ratio
SV	Science Validation
USDF	United States Data Facility
YAML	Yet Another Markup Language
photo-z	photometric redshift