



A Self-Supervised and Distillation-Guided Framework for Robust Cloud-Layer Identification in Multiwavelength Polarization Raman Lidar Networks

Weijie Zou¹, Zhenping Yin¹, Tingxin Sun², Zhichao Bu³, Yaru Dai³, Haoifei Wang⁴, Longlong Wang²,
5 Yun He^{5,6}, Shuangliang Li^{1,7}, Jiarui Cheng¹, Xuan Wang¹, Siwei Li^{1,7}, Detlef Müller¹

¹ School of Remote Sensing and Information Engineering, Wuhan University, Wuhan, China

² Hangzhou Institute of Advanced Studies, Zhejiang Normal University, Hangzhou, China

³ Meteorological Observation Center, China Meteorological Administration, Beijing, China

10 ⁴ National Satellite Meteorological Center, China Meteorological Administration, Beijing, China

⁵ State Observatory for Atmospheric Remote Sensing, Wuhan, China

⁶ School of Earth and Space Science and Technology, Wuhan University, Wuhan, China

⁷ Hubei LuoJia Laboratory, Wuhan University, Wuhan, China

Correspondence to: Zhenping Yin (zp.yin@whu.edu.cn) and Detlef Müller (dgmue@whu.edu.cn)

15 **Abstract.** Cloud-layer identification is a key prerequisite for automated ground-based lidar processing, but remains
challenging in Raman lidar networks because of incomplete near-range overlap, weak high-level cloud echoes, day–night
signal-to-noise ratio (SNR) contrasts, energy fluctuations, and aerosol–cloud ambiguity. Leveraging the China Aerosol
Raman Lidar NETwork (CARLNET), we propose a multi-channel cloud-layer identification framework that operates on
355/532/1064 nm elastic signals and 355/532 nm volume depolarization ratio, reducing dependence on absolute calibration
20 and retrieved optical products. The framework combines an anisotropic encoder–decoder architecture with a three-stage
training strategy, including self-supervised pretraining, traditional-algorithm-guided probabilistic distillation, and fine-tuning
with limited expert refinement. Strict cross-site and cross-time evaluation on held-out sites shows that, without using any
test-site labels, the framework achieves an F1 score of 0.9371 and reduces the mean absolute errors of cloud-base and cloud-
top heights to 113 m and 213 m, respectively, outperforming training without pretraining by approximately 50 m and 70 m
25 and the conventional baseline for cloud-top height by about 400 m. A labeled-data-size ablation further shows improved
label efficiency, with near-plateau performance reached at about 40 labeled days. Consistency checks against co-located
radiosonde moist-layer indications support the realism of identified high-level weak-echo clouds and the robustness of cross-
site deployment. These results demonstrate a transferable and calibration-decoupled cloud-layer identification technique for
network-scale Raman lidar processing.



30 1 Introduction

Ground-based lidar systems provide vertically resolved observations of atmospheric aerosols and clouds with high spatial resolution and continuous temporal sampling, making them important instruments for automated atmospheric profiling and long-term network observations (Campbell et al., 2002; Illingworth et al., 2007; Baars et al., 2016; Yin et al., 2020; He et al., 2024; Jing et al., 2025). Multi-wavelength Raman–polarization lidars further provide complementary information from elastic backscatter, Raman signals, and depolarization measurements, which can support aerosol/cloud discrimination, optical retrieval, and cloud-structure characterization (Sassen, 1991; Zhao et al., 2012; Thorsen et al., 2015; Alizadeh-Choobari and Gharaylou, 2017; Baars et al., 2017). However, the information content of lidar observations is strongly height dependent. In the near range, incomplete overlap and the instrument blind zone can make low-level clouds or cloud bases difficult to detect reliably. At higher altitudes, weak backscatter signals, solar background, and molecular-noise dominance make optically thin clouds, especially high-level cirrus, vulnerable to missed detections and uncertain boundary placement. These characteristics make robust cloud-layer identification a necessary but non-trivial step in automated lidar processing (Biniatoglou et al., 2018).

Cloud-layer identification is not merely a binary classification problem. In an operational lidar processing chain, cloud masks and cloud boundaries provide essential constraints for quality control, reference-layer selection, absolute calibration, attenuation correction, and subsequent aerosol/cloud optical retrievals (Klett, 1981; Fernald, 1984; Wang and Sassen, 2001; Illingworth et al., 2007; Thorsen et al., 2015; Baars et al., 2016, 2017; Yin et al., 2020). If low clouds are missed or poorly characterized in the near-overlap region, the valid retrieval range and attenuation conditions may be misinterpreted. If thin high clouds or subtle multilayer structures remain undetected, the reference region used for calibration may be contaminated, causing systematic biases in calibrated backscatter and extinction products. Conversely, aerosol layers or noise fluctuations misclassified as clouds can unnecessarily interrupt retrievals and reduce data availability. Therefore, a practical cloud-layer identification method should provide stable boundary delineation across both low-level and high-level cloud conditions, while maintaining robustness to noise, attenuation, and aerosol–cloud ambiguity.

These requirements become more demanding for atmospheric lidar networks. Compared with single-site observations, network operations involve stronger heterogeneity in instrument configuration, laser energy, receiver response, maintenance status, background light, aerosol environment, and data completeness. Such factors can alter signal-to-noise ratio (SNR), relative channel scaling, and the apparent contrast between clouds, aerosols, and clear air (Illingworth et al., 2007; Baars et al., 2016; Yin et al., 2020). As a result, methods that work well at one site or under one observing condition may not transfer directly to other sites without re-tuning. A network-scale cloud-layer identification algorithm should therefore operate on observables available before fine-grained optical inversion, reduce dependence on absolute calibration, and maintain consistent behavior under cross-site and cross-time shifts.

Cloud-layer identification in atmospheric lidar networks must therefore be evaluated under diverse and physically challenging cloud scenes rather than under simplified single-layer or high-contrast cases alone. Low-level clouds test the



ability of an algorithm to handle incomplete overlap, blind-zone effects, and strong attenuation, whereas high-level optically thin clouds test its sensitivity to weak and diffuse echoes. Multilayer scenes further require the method to preserve vertically separated structures without confusing clouds with aerosol layers or noise fluctuations (Wang and Sassen, 2001; Thorsen et al., 2015; Baars et al., 2017; Binietoglou et al., 2018). These challenges are not limited to any specific cloud type; instead, they reflect the practical observability limits and ambiguity sources of ground-based lidar measurements. Therefore, the key requirement for a transferable cloud-layer identification method is to maintain stable boundary delineation and multilayer preservation under weak echoes, attenuation, aerosol–cloud ambiguity, and cross-site observational perturbations.

65 A large body of prior work has developed rule-based and signal-processing approaches for lidar cloud-layer identification (Pal et al., 1992; Wang et al., 2018). Conventional thresholding methods are physically intuitive and computationally efficient, but their decisions are often coupled to signal amplitude, calibration constants, overlap conditions, and atmospheric background. Gradient- or change-point-based methods can capture layer transitions to some extent, but they remain sensitive to noise, attenuation, and short-term energy variability. The Value Distribution Equalization (VDE) method, inspired by

75 image histogram equalization, enhances the contrast of elastic lidar echo distributions and identifies cloud and aerosol layers based on the equalized signal structure (Zhao et al., 2014). Nevertheless, because such methods primarily rely on single-channel or empirically transformed signal characteristics, they may still underestimate cloud tops, miss weak high-level clouds, produce unstable boundaries in multilayer scenes, or generate false detections under strong aerosol loading and low-SNR conditions.

80 Recent advances in machine learning and deep learning have opened new possibilities for lidar cloud-layer identification (Cromwell and Flynn, 2019; Chisari et al., 2024, 2025). However, several limitations remain for operational deployment. First, many existing studies treat lidar time–height observations as generic two-dimensional images, without explicitly considering the physical asymmetry between the temporal and vertical dimensions (Cromwell and Flynn, 2019; Wijnands et al., 2024). Second, some approaches use retrieved backscatter or extinction coefficients as model inputs (Vainiger et al.,

85 2022). Because these retrieved quantities themselves depend on calibration assumptions, reference-layer selection, and inversion settings, using them as inputs may introduce an implicit coupling between cloud detection and downstream optical retrieval. Third, many models are trained and validated on limited sites, single regions, or within-site random splits, so the training and testing distributions remain highly similar in both site and instrument domains (Cromwell and Flynn, 2019; Fu et al., 2025). Consequently, their true cross-site transferability and deployment risk under network-scale operations are

90 difficult to assess objectively.

At network scale, the central task is therefore to infer stable cloud-layer structures from multi-channel time–height observations under cross-site shifts in instrument response and background conditions, while avoiding site-specific tuning and reducing reliance on dense manual annotations. To address this need, we propose a self-supervised and distillation-guided cloud-layer identification framework for multi-channel Raman lidar observations. The method operates directly on

95 non-absolutely calibrated elastic signals at 355, 532, and 1064 nm, together with volume depolarization ratios at 355 and 532 nm. Retrieved optical products are not used as direct model inputs, but only as auxiliary information for expert annotation



and physical interpretation. This design aims to decouple cloud-layer identification from fine-grained optical inversion and to improve robustness in automated processing chains.

The proposed framework combines a lidar-specific anisotropic encoder–decoder network with a three-stage training strategy.

100 The network explicitly accounts for the asymmetric physical meanings of the time and height axes: vertical downsampling enlarges the receptive field for multilayer structure extraction, whereas the temporal dimension is preserved to maintain short-term continuity. Training is conducted in three stages: self-supervised reconstruction pretraining to learn transferable time–height and cross-channel representations, VDE-guided probabilistic distillation to introduce physically motivated weak supervision, and fine-tuning with limited expert-refined labels to improve boundary placement. This strategy is designed to
105 improve label efficiency, suppress aerosol-induced false positives, and enhance robustness to SNR changes, energy fluctuations, and cross-site observation shifts.

We evaluate the framework using observations from the China Aerosol Raman Lidar Network (CARLNET) under strict cross-site and cross-time settings (Shao et al., 2025). Three sites representing different environmental backgrounds are held out from both pretraining and supervised training, and continuous observations from these sites are used for independent
110 testing. The evaluation includes pixel-level segmentation metrics, operation-relevant cloud-base and cloud-top height errors, qualitative analyses under complex echo conditions, and independent consistency checks with co-located radiosonde observations. Through these experiments, we assess whether the proposed method improves not only binary cloud detection accuracy, but also boundary stability, multilayer preservation, and operational applicability for lidar-network processing.

The remainder of this paper is organized as follows. Section 2 describes the CARLNET observations, lidar input channels, auxiliary optical products, manual annotation protocol, radiosonde reference data, and VDE-derived distillation labels.
115 Section 3 presents the proposed anisotropic encoder–decoder framework and the three-stage training strategy. Section 4 reports the cross-site and cross-time validation results, including label-efficiency experiments, cloud-base/cloud-top evaluation, qualitative case studies, and radiosonde-based consistency analyses. Section 5 summarizes the main conclusions, discusses limitations, and outlines future directions for operational deployment and cross-network evaluation.

120 2 Lidar observations and dataset

2.1 Ground-based lidar network and site selection

The data used in this study were collected from the China Aerosol Raman Lidar Network (CARLNET), an operational ground-based lidar network operated by the China Meteorological Administration. CARLNET consists of 47 multi-wavelength Raman–polarization lidar sites distributed across different climatic and environmental regimes in China,
125 including arid inland regions, monsoon-influenced regions, high-elevation areas, and humid coastal regions (Shao et al., 2025). And it was expanded to 79 sites in 2026, according to personal communication. This spatial coverage provides diverse observing conditions in terms of aerosol loading, humidity background, solar radiation, cloud occurrence, and



instrument operating state, making the dataset suitable for evaluating the robustness and transferability of cloud-layer identification methods under realistic network operations.

130 In this study, the CARLNET observations are used primarily for method development and cross-site validation rather than for climatological statistics. Observations from 2024 are used as the main data source for self-supervised representation learning, VDE-guided probabilistic distillation, and sampling of expert-refined labels. This design allows the model to learn from a wide range of operational conditions, including day–night background differences, signal-to-noise ratio variations, laser-energy fluctuations, aerosol–cloud coexistence, and intermittent data-quality changes. All data used for model training
135 and validation are processed under the same preprocessing protocol described below.

To rigorously assess cross-site and cross-time generalization, three sites with contrasting environmental backgrounds are selected as independent test sites: 51628 (Aksu), located in a temperate continental arid region; 53845 (Yan’an), located in a temperate monsoon-influenced region; and 58847 (Fuzhou), located in a subtropical humid coastal region. These sites differ substantially in aerosol environment, moisture conditions, background radiation, and cloud characteristics, and therefore
140 provide a high-contrast benchmark for network-scale deployment. Importantly, these three sites are excluded from both the self-supervised/distillation stages and the supervised fine-tuning stage. Unless otherwise specified, all model evaluations are conducted using continuous observations from these held-out sites in January 2025, ensuring that the test data are independent of the training data in both site and time.

The remaining 44 sites are used as training and validation sources. This split is designed to evaluate whether the proposed
145 framework can learn transferable cloud-layer representations from heterogeneous network observations and apply them to unseen sites without site-specific retuning. Figure 1 shows the spatial distribution of the CARLNET sites and the train–test split used in this study.

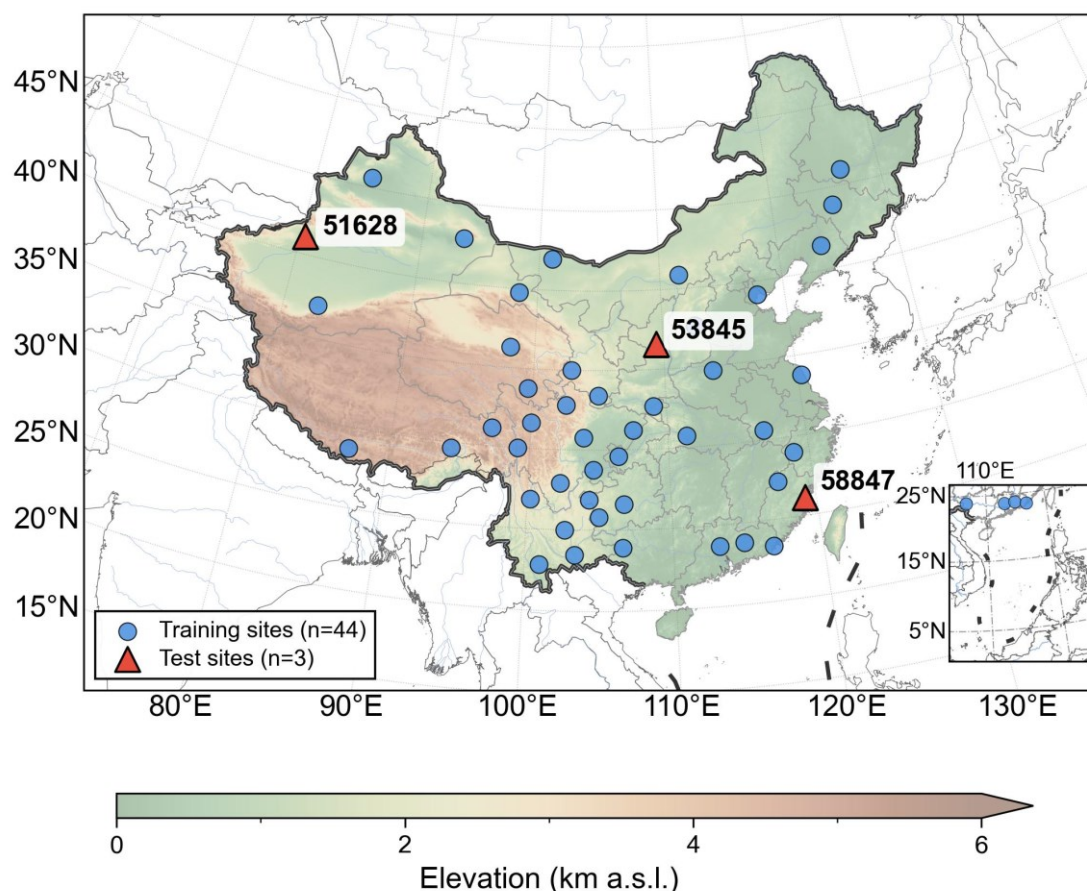


Figure 1. Spatial distribution of CARLNET sites and the schematic train–test split used in this study. Blue circles denote training/validation sites ($n = 44$), and red triangles denote the three independent test sites ($n = 3$): 51628 (Aksu), 53845 (Yan'an), and 58847 (Fuzhou). The background shows topographic elevation (km).

2.2 Lidar instrumentation and observation configuration

Ground-based lidar systems deployed in CARLNET mainly come from two vendors. Despite vendor differences, all systems are designed and integrated under unified operational technical specifications, and each site performs routine maintenance and periodic calibration following operational protocols.

The lidar used in this study is a three-wavelength polarization Raman lidar system emitting at 355 nm (pulse energy ≥ 0.6 mJ), 532 nm (≥ 1.5 mJ), and 1064 nm (≥ 1.0 mJ), with a repetition rate of 1000 Hz. And it has eight receiving channels, 355 nm parallel and vertical channels, 386 nm and 407 Raman channels, 532 nm parallel and vertical channels, and 1064 nm elastic channel. Each channel incorporates a photon counting digitizer to convert backscatter light into photon counts. The system has a vertical resolution of 15 m, a complete overlap height starting from approximately 500 m. For operational processing, raw observations are resampled to a 5 min temporal resolution to balance temporal continuity, signal stability,



and computational efficiency for automated cloud-layer identification. Details about the lidar network and lidar instrument specifications can be found in Shao et al. (2025) and Bu et al (2026).

165 Since cloud-layer identification is a prerequisite for subsequent instrument calibration and optical retrieval, we mainly use uncalibrated observables as model input features, including the elastic signals from the 355-nm elastic channel, 532-nm parallel channel, and 1064-nm parallel channel, together with linear Volume Depolarization Ratio (VDR) at 355 and 532 nm. The elastic-channel signals are preprocessed with trigger-delay correction, background subtraction, and range correction. The volume depolarization ratio is defined as

$$\text{VDR} = \text{PGR} \times \frac{P_s}{P_p}, \quad (1)$$

170 where P_s and P_p denote the elastic echo signals measured in the cross-polarized and co-polarized channels, respectively, and PGR is the polarization gain ratio provided by each site after calibration.

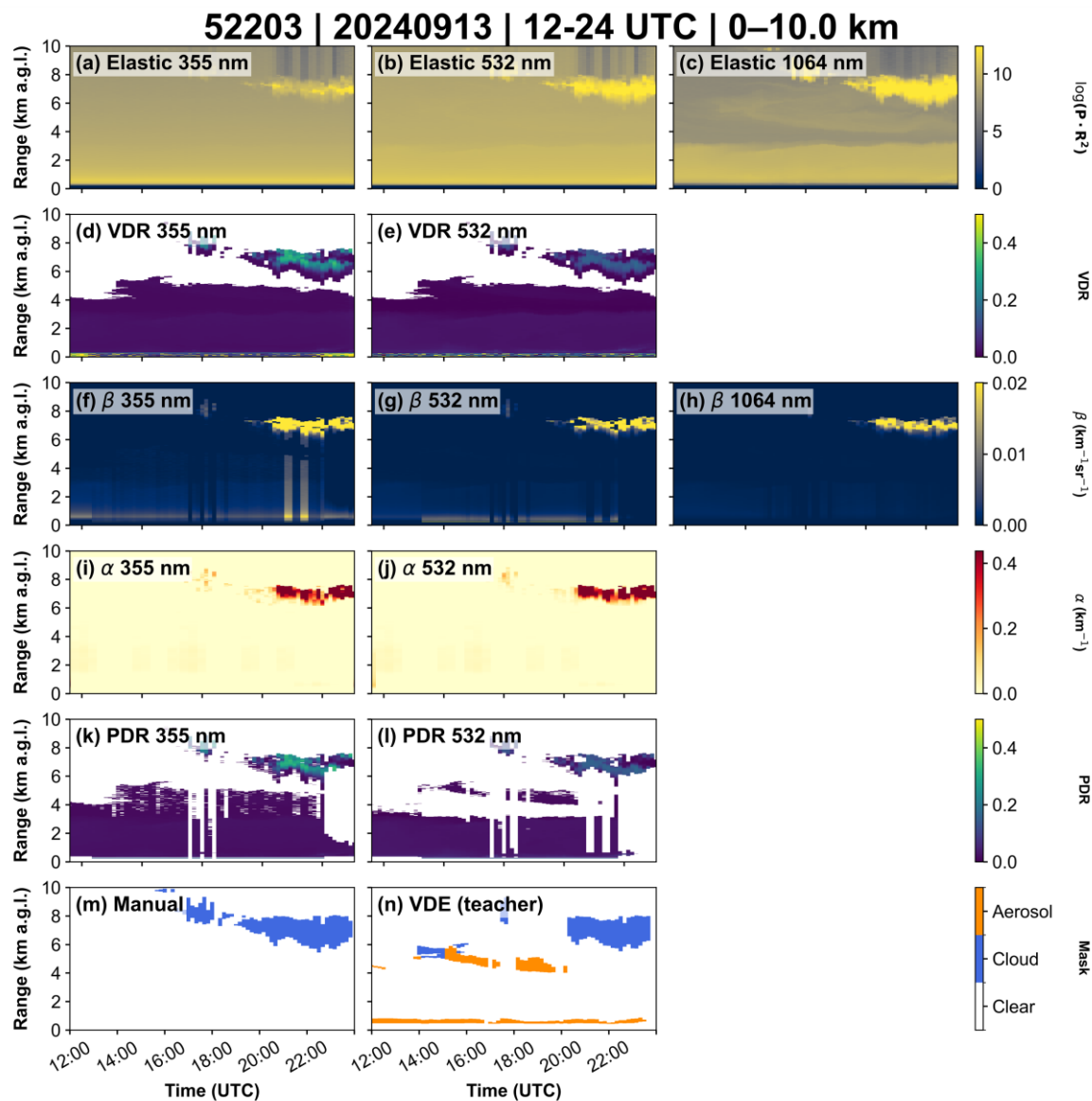


Figure 2. Example of multi-wavelength lidar observations, retrieved optical properties, and cloud-layer identification at Site 52203 on 24 December 2024 (12:00–24:00 UTC). Panels (a–c) show elastic backscatter signals at 355, 532, and 1064 nm elastic channel (range-squared corrected and displayed in logarithmic scale). Panels (d–e) show the volume depolarization ratio (VDR) at 355 and 532 nm. Panels (f–h) show retrieved aerosol/cloud backscatter coefficients β at 355, 532, and 1064 nm elastic channel. Panels (i–j) show retrieved extinction coefficients α at 355 and 532 nm. Panels (k–l) show retrieved particle depolarization ratio (PDR) at 355 and 532 nm. These retrieved optical products provide auxiliary physical constraints to enhance cloud–aerosol discrimination. Panel (m) shows the manually refined cloud/aerosol/clear-sky mask, and panel (n) shows the discretized cloud decision produced by the conventional VDE method, used for comparison and as a distillation reference. As illustrated in (n), VDE captures the major cloud structures but tends to yield persistently elevated cloud-probability levels in the lower–middle troposphere under strong aerosol loading, reflecting intrinsic limitations of empirical-threshold-driven cloud/non-cloud decisions.



2.3 Auxiliary optical products for annotation and physical interpretation

Although the proposed model is designed to operate directly on observation-level elastic and polarization quantities, expert annotation of cloud, aerosol, and clear-sky regions benefits from additional physical information. Elastic echo intensity alone can be ambiguous: strong aerosol layers may produce echoes comparable to weak clouds, low-level liquid or mixed-phase clouds may cause strong attenuation above the cloud, and high-level optically thin clouds may only appear as weak and diffuse structures. Therefore, in addition to the model input channels, we generate auxiliary optical products to support manual interpretation and to provide a physically consistent context for evaluating cloud–aerosol discrimination.

The auxiliary products include multi-wavelength backscatter coefficients β , extinction coefficients α , and particle depolarization ratio (PDR). When Raman signals have sufficient signal-to-noise ratio, these products are retrieved using Raman lidar methods (Ansmann et al., 1990). However, Raman signals can have low SNR, especially during daytime, which may lead to large fluctuations in retrieved extinction coefficients. To extend the temporal availability of optical references, when Raman retrieval is unavailable or unstable, the Fernald method is used to retrieve backscatter coefficients and to estimate extinction coefficients under an assumed lidar ratio (Fernald, 1984). This Fernald-based estimate is hypothesis-driven and generally less reliable than Raman retrieval, but it provides useful auxiliary information for interpreting cloud structure, attenuation behavior, and aerosol–cloud transitions.

It is important to emphasize that these retrieved optical products are not used as input channels in any training or inference stage of the deep learning model. They are used only for expert annotation, sample inspection, qualitative analysis, and physical interpretation. This separation is essential for the purpose of this study: the cloud-layer identification model should operate before fine-grained optical inversion in the automated processing chain, rather than depend on inversion products that themselves require cloud screening, reference-layer selection, and calibration assumptions.

Figure 2 illustrates the complementary roles of observation-level inputs and auxiliary optical products under a typical cloud–aerosol coexistence case. Elastic signals indicate scattering structures but may be ambiguous under strong aerosol loading or weak-cloud conditions. VDR and PDR provide information on particle nonsphericity and are particularly useful for identifying ice clouds and dust-like aerosol layers. Retrieved β and α further help characterize optical thickness, attenuation, and wavelength dependence. Together, these quantities provide a multi-physics background for expert annotation and for interpreting model behavior, while the actual model inputs remain limited to elastic signals and VDR.

2.4 Manual annotation and label construction

Manual annotations are produced based on a joint interpretation of multi-channel lidar observations and retrieved optical products. We stress that cloud–aerosol discrimination is not purely an “intensity saliency” problem: strong aerosol layers can in some cases yield echo intensities comparable to thin clouds, while strong attenuation in low-level water clouds or mixed-phase clouds can suppress returns above the cloud, making cloud boundaries and multilayer structures more difficult to



identify (Mylonaki et al., 2021; Floutsi et al., 2023). This is one motivation for incorporating multi-channel information and a pretraining strategy to enhance robustness.

During annotation, clouds, aerosols, and clear sky are distinguished mainly according to the following criteria:

(i) High-level cirrus: clear layered structures in elastic signals, with typically higher VDR and PDR (e.g., >0.1) (Voudouri et al., 2020).

(ii) Low-level water clouds / mixed-phase clouds: often accompanied by pronounced signal attenuation and high backscatter levels; retrieved backscatter coefficients can approach or exceed $0.01 \text{ km}^{-1}\text{sr}^{-1}$, with weak wavelength dependence (Kalesse-Los et al., 2022). For example, the cirrus layer in Figure 2m at 6–8 km exhibits both enhanced backscatter and elevated particle depolarization.

(iii) Strong aerosol layers (humid haze/urban pollution): typically characterized by lower depolarization ratios (e.g., Figure 2d–e at 0–2 km), evident spectral differences (color ratio deviating from 1), stronger wavelength dependence, and often larger lidar ratios (Floutsi et al., 2023).

(iv) Dust aerosols: usually show higher depolarization ratios but relatively lower backscatter coefficients. Their polarization signatures can partially resemble high-level cirrus, which may cause confusion if relying on volume depolarization ratio alone. Compared with cirrus, dust layers often exhibit stronger spatial coherence and more stable layered structures, with slower temporal variation in layer altitude; these properties provide auxiliary cues (Mona et al., 2012).

Manual annotation is independently performed by two experts with experience in lidar operations and research, followed by two rounds of cross-checking and sampled re-inspection to ensure consistency.

The dataset covers year-long observations from 47 sites in 2024, totaling approximately 15,000 days, with substantial spatial coverage and temporal continuity. Training data are obtained via sampling: approximately 150 days are sampled from 2024 observations of training sites, excluding the independent test sites. Test data consist of continuous observations from the three independent test sites in January 2025, totaling about 88 days, for strict cross-site and cross-time evaluation.

2.5 Radiosonde observations and cloud identification

To provide independent observational validation of cloud-layer identification results, we incorporate co-located radiosonde measurements. Radiosondes are launched daily at 00:00 and 12:00 UTC. The radiosonde station is within 100 m of the lidar site and is thus considered co-located. The radiosonde provides profiles of temperature, relative humidity, and pressure from the surface up to about 16 km, with an average ascent rate of approximately 400 m/min. The radiosonde measurements are recorded at a temporal resolution of 1 s. The corresponding vertical sampling distance depends on the instantaneous ascent velocity of the radiosonde and therefore varies with height.

Table 1. Relative-humidity thresholds used for radiosonde-based cloud identification following Zhang et al. (2016).

Altitude / km	RH
---------------	----



	Min-RH	Max-RH	Inter-RH
0 ~ 2	90% ~ 92%	93% ~ 95%	82% ~ 84%
2 ~ 6	88% ~ 90%	90% ~ 93%	78% ~ 80%
6 ~ 12	80% ~ 88%	80% ~ 90%	70% ~ 78%
> 12	75%	80%	70%

245 We apply a relative-humidity thresholding method to identify radiosonde cloud boundaries. Specifically, following the altitude-dependent RH threshold scheme proposed by Zhang et al. (2016), different minimum, maximum, and intermediate RH thresholds are used for different altitude ranges (Table 1). Based on this scheme, we adjust the order of moist-layer identification and merging steps to avoid prematurely discarding thin layers before merging. The processing steps include:

- (i) remove moist layers with base height < 120 m and cloud geometrical thickness (CGT) < 400 m;
- 250 (ii) identify moist layers using the altitude-dependent minimum RH threshold;
- (iii) merge adjacent moist layers when their separation is < 300 m;
- (iv) remove moist layers with CGT < 30 m.

It should be noted that radiosondes drift horizontally during ascent due to winds, which can be as large as 30 km. Therefore, the observations cannot be mapped to an instantaneous single-column profile at a specific time. Therefore, radiosonde-based
255 cloud identification is used only as an independent trend-level reference for external consistency checks of lidar-based cloud-layer identification, rather than as direct supervision.

2.6 VDE-based cloud-layer identification product

To improve model generalization across sites and observing conditions, we introduce a conventional cloud-layer identification approach based on elastic lidar echoes—the VDE algorithm (Zhao et al., 2014)—to construct the teacher labels
260 required for the distillation stage. This can be viewed as a single-source, experience-based information stream.

VDE is a physics-informed, rule-based method for cloud–aerosol discrimination, originally developed for cloud-layer identification from single-channel elastic lidar observations. It operates directly on elastic echoes without range-squared correction. Rather than applying a fixed intensity threshold, VDE performs a histogram equalization-like transformation on the vertical distribution of echo intensity, enhancing structural differences and highlighting potential feature layers, as shown
265 in Figure 2n.

Based on the equalized signal, VDE analyzes salient layers and combines (i) local peak characteristics and (ii) height-dependent attenuation behavior to distinguish clouds from aerosols. Typically, clouds manifest as prominent intensity peaks accompanied by rapid signal decay, especially for optically thick water clouds or low-level clouds. In contrast, aerosol layers tend to show smoother intensity variations and often lack the combined signature of strong peaks and sharp attenuation.

270 Accordingly, VDE classifies salient layers in each vertical profile: layers with distinct peaks and rapid attenuation are treated as positive samples (cloud), whereas layers with smooth variations and lacking pronounced attenuation are treated as



negative samples (aerosol or clear sky). This procedure provides an empirical discrimination mechanism with clear physical intuition.

Building on the conventional binary VDE output, Zou et al. (2024) further extended the decision from hard classification to a discretized cloud-probability representation, where different probability levels are assigned based on peak saliency and attenuation strength, thereby partially preserving uncertainty in cloud/non-cloud decisions.

As shown in Figure 2n, VDE identifies the main cloud structures but tends to generate persistently elevated cloud-probability levels in the lower–middle troposphere under strong aerosol loading, revealing inherent limitations of empirical-threshold-driven cloud–aerosol discrimination. In this study, the VDE-derived cloud-probability product is not used as the final cloud-layer identification result; instead, it is used only as weak supervision for distillation during training. Specifically, we run VDE independently on the elastic channels at 355, 532, and 1064 nm elastic channel, and then average the cloud-probability outputs across the three channels to form the final distillation soft label. This soft label guides the deep model—after self-supervised pretraining—to learn structural representations consistent with traditional physics-informed decision logic.

We emphasize that VDE is used only during the distillation stage. It is not used in final inference or performance evaluation. Its role is to provide an auxiliary constraint rooted in physical experience, rather than to replace manual annotations or the deep learning model itself.

3 Methodology

This section presents the proposed cloud-layer identification method and its training strategy. Motivated by (i) the continuous structural patterns of ground-based lidar observations in the time–height domain and (ii) the pronounced observation shifts across sites and instruments, we model multi-channel lidar measurements as a multi-channel 2D time–height array. We then develop an anisotropic encoder–decoder network together with a three-stage training framework to enable stable identification of cloud-layer structures and to improve cross-site generalization. The overall pipeline is illustrated in Figure 3. Using multi-channel time–height observations as input, the model is trained under a unified architecture in three consecutive phases: self-supervised reconstruction pretraining, distillation with VDE teacher signals, and fine-tuning with manually refined labels.

This section focuses on problem formulation, network design, and the stage-wise training strategy. Implementation details, hyperparameter settings, and additional experiments are deferred to the Appendix to ensure reproducibility.

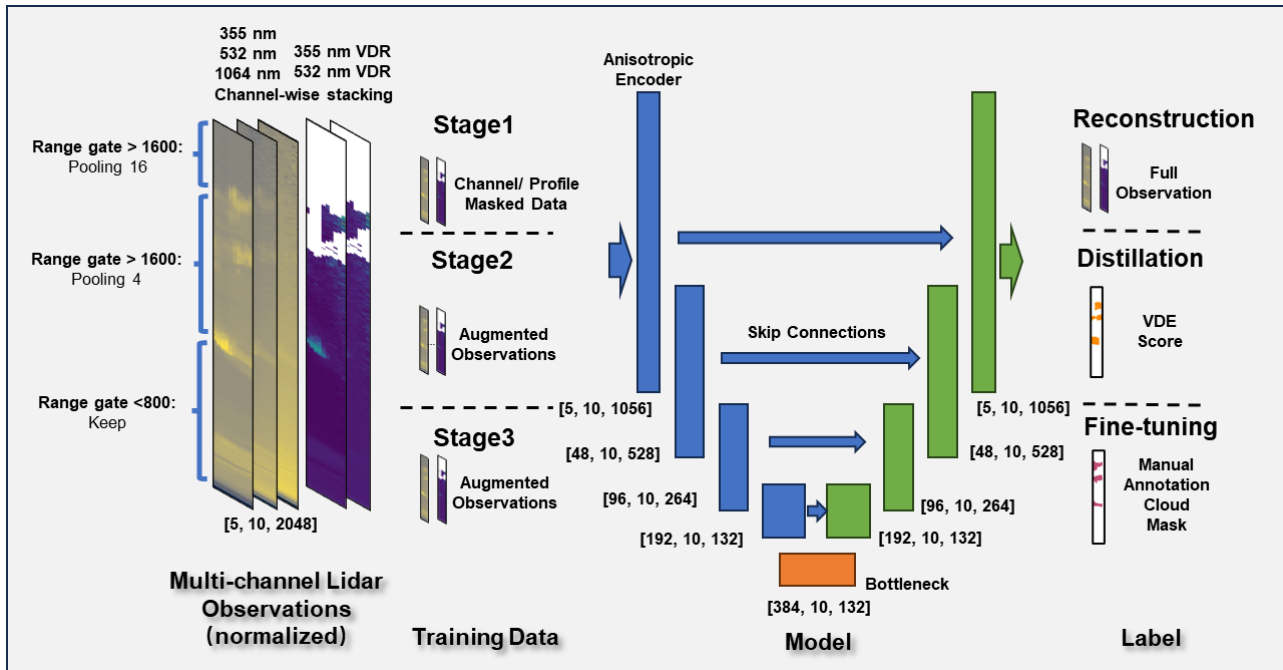


Figure 3. Schematic of the cloud-layer identification network based on multi-channel lidar observations and the three-stage training framework.

3.1 Problem formulation and input representation

Ground-based lidar actively probes the atmosphere in a time–height profiling manner, featuring high vertical resolution and strong temporal continuity. Unlike conventional 2D remote-sensing imagery, the time and height axes are physically asymmetric: the height dimension corresponds to the true vertical atmospheric structure, whereas the time dimension reflects the temporal evolution of atmospheric processes. Therefore, simply treating lidar measurements as an isotropic 2D image and modelling them uniformly is insufficient to capture the underlying physical characteristics.

In this study, multi-wavelength elastic signals (355, 532, and 1064 nm elastic channel) and VDR at 355 and 532 nm are stacked along the channel dimension to form a multi-channel time–height input tensor. The model outputs a corresponding binary cloud-layer mask, which marks echo regions associated with cloud processes.

Along the time dimension, the model does not make decisions from a single profile in isolation. Instead, it uses a finite number of adjacent profiles as a joint input. Specifically, each sample consists of 10 consecutive lidar profiles with a temporal resolution of 5 min. This design aims to incorporate local temporal continuity while avoiding the non-stationarity introduced by overly long temporal windows.

The resulting cloud-layer masks are further converted into operation-relevant geometric products, including cloud-base and cloud-top heights, for the validation analyses in Section 4.



3.2 Anisotropic encoder–decoder network architecture

To explicitly reflect the asymmetric physical meaning of time and height in ground-based lidar observations, we design an anisotropic encoder–decoder architecture (Figure 3) to fuse multi-channel measurements and extract dominant
320 spatiotemporal features for cloud-layer identification. Different from approaches that treat the time-height signal as an isotropic 2D image, our network explicitly differentiates modeling along the two dimensions.

Along the height dimension, the network performs progressive downsampling and feature compression to enlarge the vertical receptive field and enhance representation of multilayer cloud structures and weak signals at high altitudes. Along the time dimension, no downsampling is applied; instead, only limited-scale convolutions are used to extract local temporal-
325 consistency features, reducing reliance on long uninterrupted time series. This design helps preserve the continuity of cloud-layer structures while decreasing sensitivity to the completeness of temporal observations.

In the time dimension, the model takes 10 consecutive lidar profiles as a joint input to provide local temporal context for cloud-layer identification. The objective is to capture stable structural patterns on short time scales rather than to explicitly model the full temporal dynamics of cloud evolution. On the one hand, cloud-layer structures often remain coherent within
330 minutes; on the other hand, operational interruptions caused by precipitation, strong background light, or instrument-state changes are common, and a short temporal window improves effective sample utilization and practical applicability.

Between the encoder and decoder, skip connections perform joint fusion across scales and channels, propagating wavelength-/polarization-specific cues from the encoder to the decoder for recovering fine-grained cloud boundaries (Ronneberger et al., 2015; Oktay et al., 2018). In addition, we introduce gated skip connections (Jiang et al., 2021) to
335 adaptively weight the transferred features, thereby regulating the contribution of multi-scale and multi-channel information under varying SNR and cloud–aerosol ambiguity. This design preserves the multi-scale representation strength of U-Net while explicitly supporting multi-channel feature fusion tailored for network-scale lidar operations, where high vertical resolution, local temporal continuity, and cross-site observation shifts must be handled simultaneously.

3.3 Three-stage training strategy

340 To improve stability and generalization under limited manual annotations, we adopt a three-stage training strategy (Figure 3). Stage 1 employs a self-supervised reconstruction task based on channel-wise and profile-wise masking, guiding the model to learn spatiotemporally coherent structures and cross-channel consistency from lidar echoes without any manual labels (Vincent et al., 2008; He et al., 2021).

Stage 2 performs VDE-guided probabilistic distillation to regularize the cloud/non-cloud decision boundary learned in Stage
345 1. Specifically, we use discretized cloud-probability levels produced by the conventional VDE algorithm (Zhao et al., 2014; Zou et al., 2024) as soft targets, and optimize the segmentation head on top of the self-supervised representation. Compared with hard labels, these probability levels provide an uncertainty-aware training signal that stabilizes learning in ambiguous conditions, such as weak echoes, strong background noise, and aerosol–cloud overlap. In addition, we explicitly include



aerosol-labeled regions as negative constraints to penalize aerosol-induced false positives and prevent the model from drifting toward overly permissive cloud decisions. Overall, Stage 2 improves transferability by injecting a robust, label-efficient regularization signal while limiting dependence on dense manual annotations.

Stage 3 fine-tunes the model using a small set of manually refined cloud-layer masks. This stage further improves boundary delineation accuracy and cross-site generalization while preserving the stable structural representations learned in the previous stages.

It should be noted that the multi-vendor and multi-site observation configurations in CARLNET inevitably introduce instrument-parameter differences and energy fluctuations in real operations. Rather than treating this heterogeneity as noise to be avoided, we leverage it as a realistic condition for evaluating robustness under varying instrument settings. By incorporating cross-site and cross-time data in the self-supervised and distillation stages, the model is encouraged to learn structural representations that are insensitive to scale variations and operational fluctuations, thereby improving generalization to unseen sites and different instrument conditions.

3.4 Data augmentation and robustness design

In real-world operations, ground-based lidar observations are inevitably affected by non-ideal factors, including cross-vendor instrument differences, laser energy fluctuations, background-noise variations, and signal perturbations under complex environmental conditions. These factors are pervasive at network scale and constitute major sources that limit the cross-site generalization of deep learning models.

In this study, raw lidar signals are resampled to a 5 min temporal resolution in operational processing to balance temporal continuity, signal stability, and computational efficiency for automated cloud-layer identification.

However, the effective time resolution may still differ across sites due to observing conditions, quality control, or missing data. To address this, we introduce a set of lidar-oriented data augmentation strategies during training. Through explicit sampling and perturbation mechanisms, these augmentations cover variations in time scales and observing conditions, thereby improving robustness to time-resolution shifts.

The adopted augmentations operate primarily at the input-observation level, aiming to emulate realistic conditions such as energy-scale variations, noise perturbations, local temporal gaps, and incomplete channel information (Xie et al., 2020). These operations encourage the model to focus on structural patterns in the time–height domain, rather than depending on specific amplitude scales or fixed sampling intervals, thus reducing sensitivity to configuration differences.

These augmentation strategies are applied in different forms across the three training stages, and play particularly important roles in self-supervised pretraining and VDE-based distillation: the former facilitates learning stable time–height structural representations, whereas the latter further suppresses identification uncertainty induced by noise and scale changes under a physically motivated constraint.



We emphasize that all data augmentations are used only during training and are not applied during inference or evaluation. Moreover, all augmentations are designed to be physically plausible for lidar observations and do not introduce any additional prior label information. For reproducibility, the augmentation types, parameter ranges, and stage-wise usage will be reported in the Appendix.

385 Once trained and validated, the model can be applied with the same preprocessing and inference protocol to generate consistent cloud-layer masks, cloud-base heights, and cloud-top heights for automated CARLNET processing. These products provide a basis for downstream calibration, quality control, and future long-term cloud-structure analyses, but the present study focuses on method validation and cross-site transferability.

4 Experiments and validation

390 This section evaluates the proposed framework as an automated cloud-layer identification technique for network-scale Raman lidar processing. The validation is designed to test whether the method can maintain stable performance under cross-site differences in instrument response, background conditions, aerosol environment, and cloud structure. We assess the framework progressively in terms of label efficiency, pixel-level identification accuracy, operation-relevant geometric products including cloud-base and cloud-top heights, and external consistency with independent radiosonde observations.

395 Unless otherwise specified, all evaluations are conducted using continuous observations from three held-out CARLNET sites (51628, 53845, and 58847) in January 2025. These sites are excluded from both pretraining and supervised training, thereby ensuring a strict cross-site and cross-time evaluation setting. Any experiment that introduces a small amount of labeled data from the test sites is explicitly treated as an adaptation ablation rather than as part of the main validation protocol. The VDE method is used as the rule-based baseline throughout this section.

400 4.1 Data-size ablation study

To evaluate how different training strategies behave under limited manual annotations, we design an ablation study with respect to the amount of labeled training data. The x-axis indicates the number of consecutive observation days used for training, and the y-axis reports the F1 score on the independent test set. The definition and computation of the F1 score are provided in Appendix A. The test set is fixed to continuous observations in January 2025 from the three test sites, ensuring

405 that training and testing are independent in both site and time.

We compare four strategies: No Pretrain, Stage-1 self-supervised pretraining only, Stage-2 distillation pretraining only, and joint Stage-1+2 pretraining. Figure 4 summarizes the performance under different training-data scales. The results show that under small-sample conditions (10 training days), pretraining yields a clear performance boost, and the joint pretraining strategy provides the most pronounced gain over training from scratch. As the training-data scale increases, all methods

410 improve but exhibit diminishing marginal returns, approaching a stable plateau at around 40 days.



This experiment indicates that self-supervised and distillation pretraining significantly improves data utilization efficiency, enabling better learning performance with fewer manual annotations and encouraging the model to learn structural representations that are more robust to cross-site differences and operational perturbations at an earlier stage. Considering the trade-off between plateau performance and training cost, the subsequent experiments adopt a training scale of approximately 40 consecutive days.

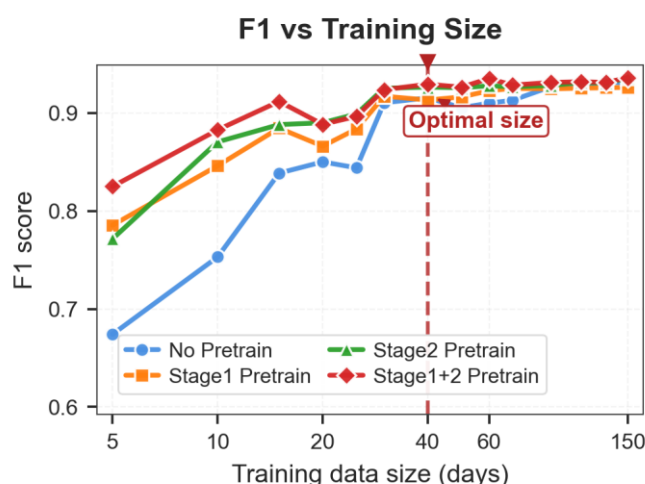


Figure 4. Comparison of cloud-layer identification performance (F1 score) under different training-data scales. The red dashed line and annotation indicate the selected training-data size of 40 days, where the performance approaches a plateau and marginal gains become limited.

Overall, the results suggest that pretraining not only reduces reliance on large-scale manual annotations for cloud-layer identification, but also provides a more controllable training cost for cross-site deployment.

4. 2 Cross-site overall performance

4.2.1 Overall performance and clarification of test-site usage

Table 2 reports the overall performance under different training strategies and different usages of the test-site data. Metrics include intersection over union (IoU), F1 score, precision, and recall. All results are computed over the full test set. We apply an early-stopping strategy to mitigate overfitting: training is terminated when the validation performance does not improve for a number of consecutive epochs.

In Table 2, “Test sites used in pretraining” indicates whether the test-site data are included during pretraining, and “Test-site labels used in training” indicates whether a small amount of labeled data is additionally introduced in the training stage for each of the three test sites (2 days per site, sampled from 2024 and independent of the test month). This setting is designed to emulate a practical scenario where a newly deployed site provides only limited annotations for adaptation. “Pretrain (no test



sites) + train (no test sites)” corresponds to the strictest configuration: the test sites are used neither in pretraining nor in training, reflecting the zero-prior generalization capability for truly unseen network sites.

435 As shown in Table 2, pretraining yields stable gains in cross-site evaluation. Under the setting where no labeled data from test sites is introduced during training, simulating newly added operational sites, the pretrained model achieves $F1 = 0.9371$, higher than 0.9281 obtained without pretraining. Measured by inconsistency ($1 - F1$), pretraining reduces it from 0.0719 to 0.0629, corresponding to a 12.5 % relative reduction. Notably, introducing a small amount of labeled data from the test sites during training does not further improve performance substantially ($F1 = 0.9347$), suggesting that the major benefit primarily
440 comes from transferable structural representations acquired during pretraining rather than from supervised fitting to specific test sites.

Compared with the deep learning models, VDE performs markedly worse under network-scale observations, particularly in IoU and F1 score, highlighting the limitation of rule-based methods under complex cloud structures and strong aerosol
445 backgrounds.

Table 2. Cloud-layer identification performance under different training strategies and different usages of test-site data.

Strategy	Pretrain	Test (Pretrain)	Sites Test (Training)	Sites IoU	F1 Score	Precision	Recall
No Pretrain (No Test Sites)	X	-	X	0.866	0.928	0.910	0.947
No Pretrain (With Test Sites)	X	-	✓	0.868	0.929	0.927	0.931
Pretrain (No Test Sites) + Train (No Test Sites)	✓	X	X	0.882	0.937	0.926	0.949
Pretrain (No Test Sites) + Train (With Test Sites)	✓	X	✓	0.878	0.935	0.924	0.945
VDE	-	-	-	0.295	0.404	0.429	0.539

Overall, the results indicate that pretraining provides consistent improvements under strict cross-site and cross-time settings, and the strictest configuration performs no worse than configurations that include test-site information, supporting that the performance gain does not originate from site-specific fitting. This suggests that pretraining primarily enhances the model’s
450 general representation of cloud-layer structures, offering a more reliable foundation for operational scaling.

4.2.2 Operation-relevant evaluation: cloud-base and cloud-top heights

The operational value of a cloud-layer mask is not only reflected by pixel-level classification accuracy, but also by its capability to characterize cloud geometrical structure—most importantly, the stable estimation of cloud-base height and cloud-top height. To assess method differences from an application-oriented perspective, we convert predicted cloud-layer
455 masks into cloud-base and cloud-top height products and compare them with manually refined annotations.



Figure 5 shows the correspondence between estimates from different methods and the manual ground truth (GT), and Table 3 summarizes quantitative statistics including the detection rate and the errors of cloud-base/cloud-top heights. For cloud-base estimation, the pretrained model achieves a mean absolute error (MAE) of 113 m, reducing the MAE by roughly 50 m compared with the no-pretrain model (162 m). For cloud-top estimation, the pretrained model yields an MAE of 213 m, improving by about 70 m over the no-pretrain model (283 m). Relative to VDE, the improvement is particularly pronounced for cloud-top height: the cloud-top MAE decreases from 622 m to 213 m, i.e., by nearly 400 m.

Beyond geometric errors, Table 3 also reports the detection rate, defined as the fraction of samples in which a cloud process is detected. The pretrained and no-pretrain models show nearly identical detection rates (94.81 % vs. 94.85 %), indicating comparable overall capability in detecting cloud processes; importantly, pretraining does not improve performance simply by detecting more clouds. In contrast, VDE exhibits a lower detection rate (64.07 %), reflecting a higher tendency to miss weak-echo clouds or complex layered structures at network scale.

A plausible explanation for VDE's lower detection rate can be inferred from Figure 5c and f. VDE tends to overestimate cloud-base height while underestimating cloud-top height, which can narrow or even eliminate marginal cloud layers in the derived vertical extent. The cloud-top underestimation is consistent with its boundary determination strategy, whereas the elevated cloud base is likely related to the smoothing/denoising steps that suppress weak echoes near cloud boundaries. Moreover, VDE shows more abnormal outliers in the geometric errors, suggesting reduced robustness under low-SNR conditions and complex multilayer scenes at network scale, thereby increasing the chance of missing clouds altogether. To further substantiate these interpretations with an independent reference, Section 4 introduces radiosonde observations for an objective evaluation.

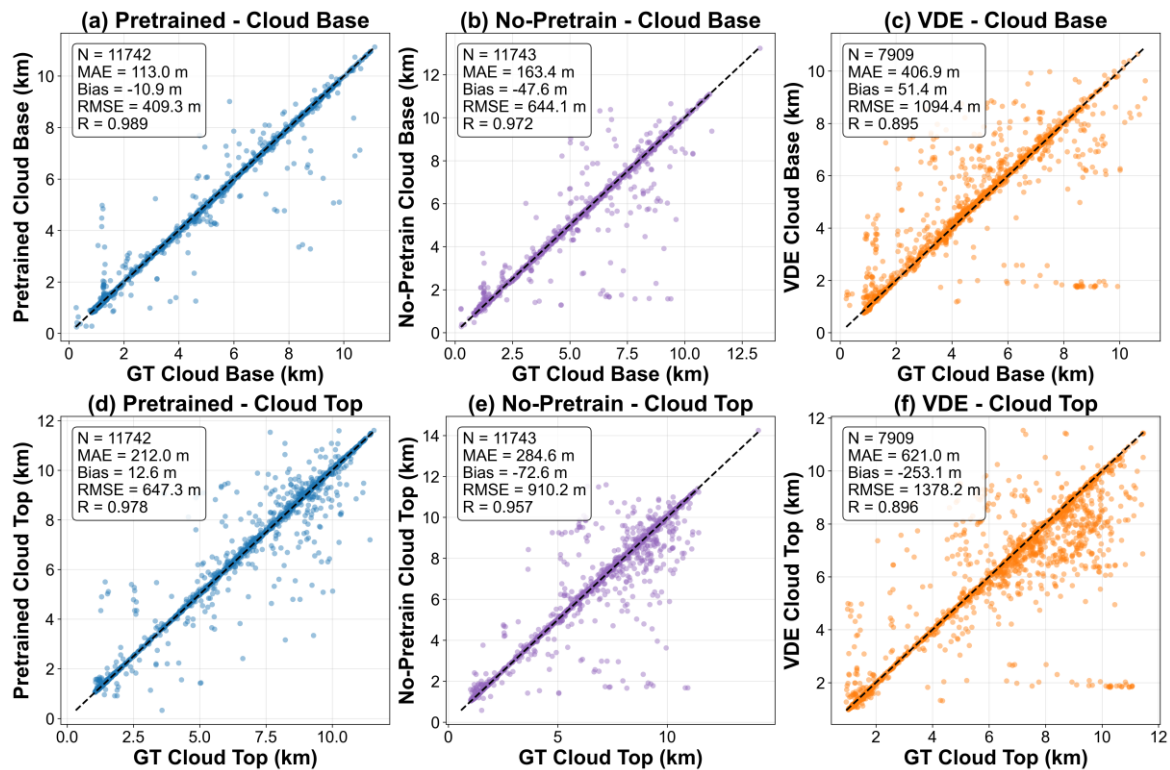


Figure 5. Cloud-base and cloud-top height estimates compared with the manual ground truth (GT). Panels (a–c): cloud-base; panels (d–f): cloud-top.

Taken together, the results suggest that the advantage of pretraining mainly lies in more accurate delineation of cloud boundaries and geometrical structure, rather than in the binary decision of “cloud present or not.” The reduced cloud-top error and bias indicate stronger structure preservation under weak echoes, multilayer conditions, and attenuation, which directly benefits lidar calibration, quality control, and cloud-structure characterization in operational applications.

Table 3. Quantitative comparison of detection rate and cloud-base/cloud-top height errors across methods. MAE denotes mean absolute error.

Method	Detection Rate (%)	Base MAE (m)	Base Bias(m)	Top MAE (m)	Top Bias (m)
Pretrained	94.8	113.3	-10.0	212.8	13.4
No-Pretrain	94.9	162.4	-46.9	283.1	-70.9
VDE	64.1	407.4	56.5	621.6	-245.2



4.2.3 Qualitative examples under complex conditions

485 To further illustrate the role of pretraining in structural stability, Figure 6 provides qualitative examples comparing the pretrained and no-pretrain models under complex echo conditions. The pretrained model produces cloud-layer masks that are closer to manual annotations in terms of weak-echo cloud layers, multilayer structures, and boundary continuity, and also shows better temporal coherence. By contrast, the no-pretrain model more often produces layer breaks, misses, or localized false alarms over strong aerosol layers.

490 Specifically, Figure 6a shows a case at Fuzhou (58847). During 10:00–18:00 UTC, a weak aerosol layer exists above the boundary layer. The no-pretrain model persistently misclassifies this weakly scattering layer as cloud (black box), resulting in a non-physical continuous cloud coverage over a long period. The pretrained model suppresses such aerosol-induced false alarms more effectively and yields more coherent mask boundaries. Figure 6b shows a case at Aksu (51628). Around 12:00 UTC, a short-term energy fluctuation occurs; the no-pretrain model misclassifies locally enhanced noise as cloud, as marked

495 by the black box, forming isolated “pseudo-cloud” fragments that are inconsistent with surrounding structures. In comparison, the pretrained model maintains a more stable decision without corresponding false alarms.

These examples support that pretraining improves structural stability rather than providing incidental metric gains: it produces more continuous boundaries, more reliable multilayer preservation, and stronger suppression under strong aerosol backgrounds. Together with the detection-rate statistics, the advantage of deep learning at network scale is mainly attributed

500 to more complete depiction of complex cloud structures, rather than merely identifying “strong-echo clouds.”

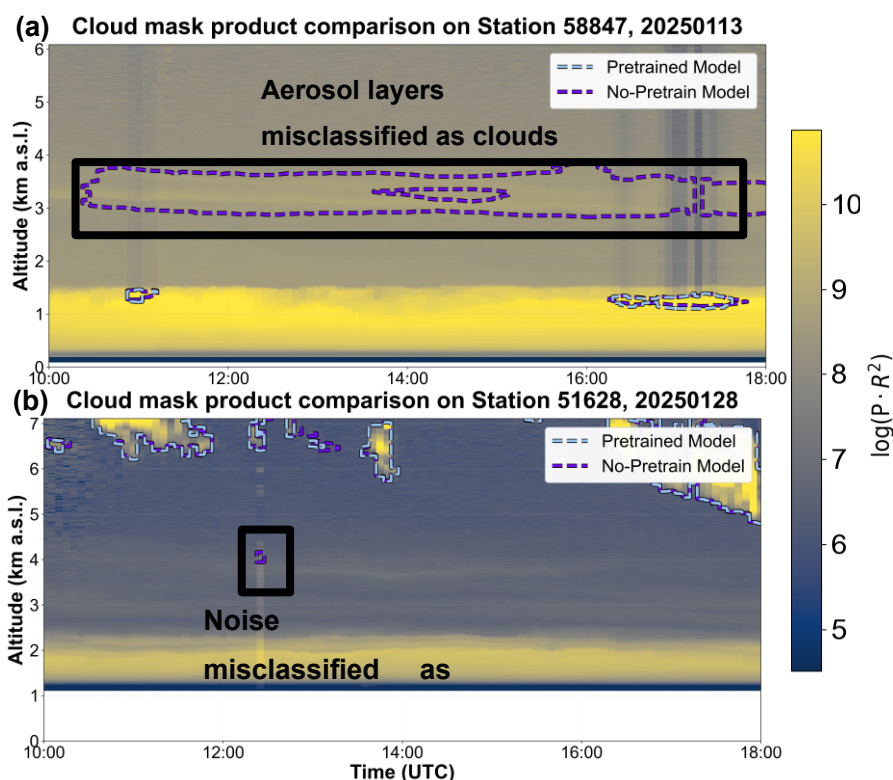


Figure 6. Qualitative comparison of cloud-layer masks produced by pretrained and no-pretrain models under complex echo conditions. The background shows the 1064 nm elastic channel (10:00–18:00 UTC; 0–6 km). The blue dashed contour denotes the pretrained model, and the green dashed contour denotes the no-pretrain model. Black boxes highlight representative regions with substantial differences. Panel a shows Fuzhou (58847, 2025-01-28), and panel b shows Aksu (51628, 2025-01-13).

4.3 Independent validation with radiosonde observations

To further assess cloud-layer identification results using an independent observing system, we incorporate co-located radiosonde data and compare the deep learning method (DL) with the conventional VDE approach.

Given the time-accumulating nature of radiosonde observations, we take the union of lidar-based cloud-layer masks within ± 1 h around each radiosonde launch time and match it with radiosonde-identified cloud layers to reduce spatiotemporal mismatch. Figures 7 and 8 show the results for Aksu (51628) and Fuzhou (58847) in January 2025.

For Aksu (Figure 7), DL and VDE are generally consistent in the temporal and vertical evolution of major cloud processes, while both often yield lower cloud bases and higher cloud tops compared with radiosondes. This systematic discrepancy mainly arises from differences in observing principles: lidar provides higher vertical resolution and sensitivity to weak echoes, whereas radiosonde moist-layer identification depends strongly on thresholding and the vertically accumulated ascent process. On this basis, DL tends to better preserve the continuity of high-level weak-echo structures (Figure 7b), with fewer temporal breaks and fewer missed thin-cirrus segments than VDE, which helps mitigate cloud-top underestimation for



tenuous upper-level layers. DL also provides more stable identification for low-level clouds, showing better agreement with radiosondes (Figure 7c).

520 For Fuzhou (Figure 8), high-level cirrus coexists with persistent low-level clouds, creating a challenging scene in which weak upper-level echoes can be intermittently obscured or fragmented. DL maintains a more continuous representation of high-level cirrus over extended periods (Figure 8b), showing better consistency with the radiosonde moist-layer intervals and exhibiting fewer intermittent gaps and missed segments than VDE at these altitudes. This improvement is particularly relevant for thin, weak-echo cirrus layers, for which rule-based methods are more prone to discontinuous detection. For low-
525 level clouds (Figure 8c), both DL and VDE capture the main structures reasonably well, and the differences are comparatively limited.

Based on the above trend-level comparison, we further quantify how well multilayer structures identified by radiosondes are preserved in the DL and VDE results, as a complementary assessment of multilayer completeness. Among 20 events in which radiosondes indicate a second cloud layer, DL identifies 17 (85.0 %), whereas VDE identifies 15 (75.0 %). For 85
530 first-layer cloud events, the identification ratios are similar (DL: 67/85, 78.8 %; VDE: 66/85, 77.6 %). This suggests that for the second layer, which is more likely to correspond to high-level or weak-echo structures, DL shows higher structural consistency than VDE, whereas for first-layer clouds dominated by strong echoes at low levels, the difference is limited. We stress that this statistic is used only to characterize structure preservation under radiosonde constraints and should not be interpreted as an absolute accuracy metric.

535 In summary, independent radiosonde-based comparisons indicate that DL and VDE are broadly consistent in the temporal evolution of major cloud processes. Under the systematic offsets expected from differing observing principles, DL exhibits clearer advantages for high-level, weak-echo layers, typical of thin cirrus, and for multilayer scenes, by preserving layer continuity and reducing intermittent misses relative to VDE. This behavior aligns with the reduced cloud-top errors and the qualitative examples under complex echoes, and supports the interpretation that pretraining mainly improves structural
540 stability and physical plausibility, rather than simply increasing cloud detections or yielding isolated metric gains.

We note that radiosonde-based cloud identification here is derived from altitude-dependent relative-humidity thresholds and is inherently uncertain due to the threshold scheme, moist-layer merging choices, and horizontal drift during ascent. Therefore, radiosonde results are used as an independent structural reference to assess consistency in layer altitude and multilayer occurrence at the trend level, rather than as a strict quantitative ground truth.

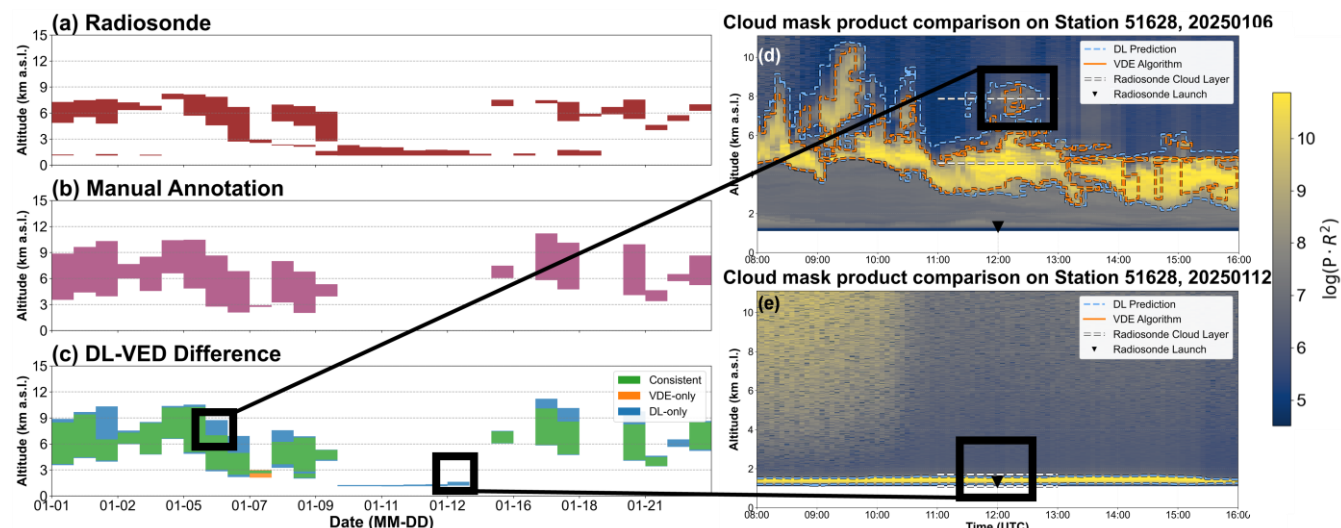


Figure 7. Consistency analysis of cloud-layer identification at Aksu (51628, January 2025) under radiosonde constraints. Panel a shows the radiosonde moist-layer cloud structure; panel b shows the manually refined labels; panel c shows the difference map between the deep learning (DL) algorithm and the VDE method (green: both; blue: DL only; orange: VDE only); panels d and e show zoom-in examples over representative periods, with the 1064 nm elastic channel as the background and overlays of the DL/VDE products and radiosonde constraints.

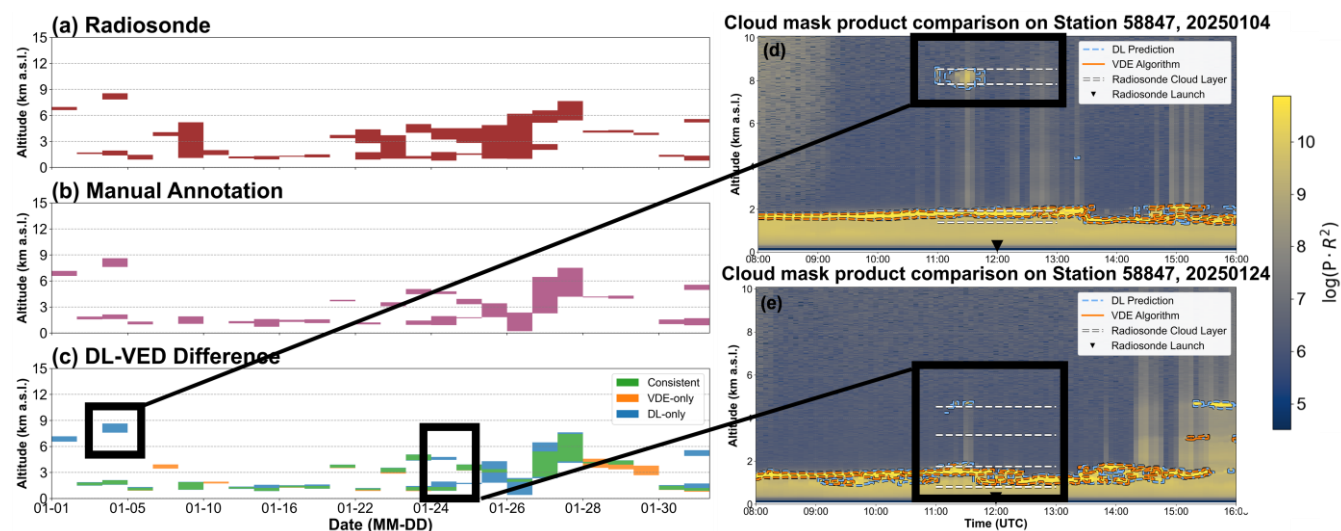


Figure 8. Same as Figure 7, but for Fuzhou (58847, January 2025).

5 Conclusion and Discussion

This study developed and validated a calibration-decoupled, multi-channel cloud-layer identification framework for automated Raman lidar network processing. The framework operates directly on observation-level elastic signals at 355, 532, and 1064 nm and volume depolarization ratios at 355 and 532 nm, without relying on fine-grained optical retrieval products



as model inputs. By combining a lidar-specific anisotropic encoder–decoder architecture with self-supervised pretraining, VDE-guided probabilistic distillation, and fine-tuning with limited expert-refined labels, the proposed method is designed to improve cloud-boundary stability, multilayer preservation, and cross-site transferability under realistic operational perturbations.

The framework was evaluated under a strict cross-site and cross-time setting using continuous observations from three held-out CARLNET sites that were excluded from both pretraining and supervised training. The results show that pretraining improves pixel-level cloud-layer identification without introducing any labeled data from the test sites. Compared with training from scratch, the F1 score increases from 0.9281 to 0.9371, corresponding to a 12.5 % relative reduction in inconsistency ($1 - F1$). The labeled-data-size ablation further demonstrates improved label efficiency, with the model approaching a performance plateau using a limited number of manually refined training days. These results indicate that the main benefit of pretraining is not site-specific fitting, but the acquisition of transferable time–height and cross-channel representations.

From an operational perspective, the improvement is more evident in geometry-related products. Compared with training from scratch, the pretrained model reduces the mean absolute error of cloud-base height by approximately 50 m and that of cloud-top height by approximately 70 m. Relative to the conventional VDE baseline, the reduction in cloud-top error is particularly pronounced, with cloud-top MAE decreasing from 622 m to 213 m. The detection rates of the pretrained and no-pretrain models remain nearly identical, indicating that the performance gain mainly arises from more accurate boundary delineation and multilayer structure preservation rather than from simply detecting more cloud pixels. Qualitative cases further show that the pretrained framework suppresses non-physical cloud detections over strong aerosol layers and reduces false alarms caused by short-term energy fluctuations.

Independent radiosonde-based comparisons provide an additional consistency check from a different observing system. Although lidar and radiosonde cloud indications are not expected to match exactly because of differences in observing principles, humidity-threshold uncertainty, and balloon drift, the proposed framework shows better preservation of second-layer structures than VDE under radiosonde constraints. This advantage is particularly relevant for high-level weak echoes and multilayer scenes, where rule-based methods are more prone to intermittent misses and cloud-top underestimation. For first-layer clouds dominated by stronger low-level echoes, the difference between the proposed method and VDE is smaller, suggesting that the primary practical value of the proposed framework lies in improving structural completeness under weak-signal and multilayer conditions.

Several limitations should be noted. First, the present framework focuses on cloud-layer identification and does not explicitly separate precipitation from cloud layers. Scenes affected by precipitation, virga, or strongly attenuating hydrometeors may require additional physical constraints or task-specific labels. Second, although the model is designed to reduce dependence on absolute calibration, its inputs remain affected by near-range overlap, daytime background radiation, signal saturation, and severe attenuation by optically thick low clouds. These factors may still limit the observability of certain cloud regions, especially near the blind zone and above persistent low clouds. Third, the model uses a short temporal context window and is



not designed to depend on long temporal averaging; however, its behavior under highly irregular sampling, sparse observations, and prolonged data gaps still requires further systematic evaluation. Fourth, although the current validation uses strict held-out sites within CARLNET, broader tests on external lidar networks are needed to assess cross-instrument and cross-platform generalizability.

595 Future work will focus on improving inference-time robustness and engineering readiness under challenging operational conditions, including strong daytime background, frequent low-cloud attenuation, irregular sampling, and data gaps. The framework will also be extended to external ground-based lidar networks and compared with spaceborne lidar observations to further evaluate cross-platform consistency. In addition, the resulting cloud-layer products can provide a unified basis for downstream calibration, quality control, aerosol/cloud optical retrieval, and future long-term cloud-structure analyses. The
600 training and inference framework, together with pretrained models, will be made available where possible to support reproducible evaluation and broader deployment.

Appendix A

This appendix summarizes the training protocol of the proposed cloud-layer identification framework, including
605 preprocessing, augmentation, model configuration, optimization, and evaluation settings. Unless explicitly stated, the same preprocessing and augmentation pipeline is applied across all stages to ensure consistent learning signals and to facilitate reproducibility.

A1. Preprocessing of multi-channel time–height observations

We distinguish preprocessing vertical compression (for computational efficiency) from architectural anisotropy, where
610 downsampling is applied only along the height axis inside the network.

All lidar channels are first normalized using a robust median absolute deviation (MAD) scheme to mitigate cross-site amplitude shifts and outlier contamination. For the photon-counting elastic channels (1064 nm elastic channel, 355 nm and 532 nm parallel channels), we further apply range-squared correction followed by logarithmic compression in order to reduce dynamic range while preserving structural contrast, i.e.,

$$615 \quad x \leftarrow \log(1 + x \cdot R^2), \quad (2)$$

where R denotes range. In contrast, the polarization-derived channels, i.e., VDR355 and VDR532, are kept in their native form and are not subject to range-squared correction, since they already represent a physically normalized ratio quantity.

To reduce computational cost while preserving the near-surface and boundary-layer structures relevant to operational usage, we compress the vertical axis using a height-dependent strategy that retains full resolution in the lower troposphere and
620 progressively coarsens the upper levels. Specifically, the 0–12 km range is kept at the original resolution, 12–24 km is downsampled by a factor of 4, and 24–30 km is downsampled by a factor of 8. Starting from 2048 range bins with 15 m



spacing, corresponding to approximately 30 km, this compression yields an effective vertical dimension of approximately 512 bins. The same vertical compression is used in all training stages and during inference.

A2. Data augmentation for operational robustness

625 Operational lidar observations are affected by instrument-to-instrument differences, laser energy fluctuations, background noise variability, and occasional data irregularities. To encourage the model to learn structural cues rather than relying on absolute amplitude scales or specific sampling patterns, we adopt a strong but physically plausible augmentation policy throughout training. Concretely, we apply (i) random height cropping with probability 0.3 using a crop ratio sampled from [0.7, 0.9], (ii) additive Gaussian noise with probability 0.3 and standard deviation 0.05, (iii) multiplicative signal scaling with
630 probability 0.3 using a factor sampled from [0.9, 1.1], and (iv) spatial dropout with probability 0.2 that randomly zeroes 10 % of pixels. These perturbations are applied only to the input observations and do not introduce any additional label information.

A3. Unified model configuration

A single unified anisotropic encoder–decoder network is used across all experiments. The input is a five-channel time–height
635 tensor comprising the 1064 nm elastic channel, the 355 nm parallel-polarized elastic channel, the 532 nm parallel-polarized elastic channel, and the volume depolarization ratios at 355 nm and 532 nm. The segmentation head outputs a one-channel cloud-layer mask (cloud vs. non-cloud). The model uses a base width of 48 channels with dropout rate 0.1 and a vertical convolution kernel size $k_z = 5$. The overall parameter count is approximately 8.5M. A reconstruction head is enabled only in Stage 1 (see below) to reconstruct all five channels; it is removed in Stage 2 and Stage 3 when training focuses on cloud-
640 layer identification.

A4. Stage-wise objectives

Training follows the three-stage strategy described in the main text.

Stage 1 (self-supervised representation learning). The model is pretrained without manual annotations using a multi-task reconstruction objective designed to capture spatiotemporal continuity and cross-channel consistency in time–height echoes.
645 The objective combines: (i) a full-channel reconstruction loss L_{recon} implemented as MSE over all channels, (ii) a channel-masking prediction loss L_{ch} where 1–2 channels are randomly masked and reconstructed, and (iii) a temporal-frame prediction loss L_t where one time frame is masked and predicted from its neighbors. This stage encourages a transferable representation that is less sensitive to site-specific scaling and short-term perturbations.

Stage 2 (weakly supervised distillation). Starting from the Stage-1 initialization, we inject physical heuristics via distillation
650 using discretized VDE-derived cloud probability as a teacher signal. We employ a weighted BCE-with-logits loss with a positive class weight $\text{pos_weight} = 19.0$ (reflecting an average cloud fraction of approximately 5%). Two additional terms are incorporated to stabilize training under operational artifacts: regions with low teacher cloud score (≤ 0.4) receive an aerosol-suppression penalty (weight 3.0) to reduce aerosol-induced false positives; and the near-range blind zone (the first 30 range bins, approximately 450 m) is strongly penalized (weight 15.0) to discourage non-physical cloud predictions in the
655 blind region.



Stage 3 (fine-tuning with manual labels). Finally, the model is fine-tuned using a limited amount of expert manual annotations with a standard BCE-with-logits loss and $\text{pos_weight} = 19.0$. This stage refines cloud boundary placement while preserving the robust structural representation learned in Stages 1–2.

A5. Optimization schedule and freezing strategy

660 All stages are optimized using AdamW (Loshchilov and Hutter, 2019). Stage 1 uses a learning rate of 5×10^{-5} with weight decay 1×10^{-4} , batch size 8, and 20 epochs. Stage 2 is trained in two consecutive phases to reduce overfitting to weak supervision: Phase 2A runs for 15 epochs with learning rate 5×10^{-5} and weight decay 1×10^{-4} while freezing the first two encoder blocks; Phase 2B runs for 15 epochs with learning rate 2×10^{-5} and weight decay 1×10^{-5} while freezing the entire encoder (decoder-only adaptation). Stage 3 fine-tuning runs for 40 epochs using learning rate 1×10^{-5} , weight decay
665 5×10^{-4} , batch size 4, and a cosine annealing schedule with $\eta_{\min} = 10^{-6}$ while freezing the encoder and bottleneck to focus adaptation on the decoder and prediction head. Early stopping is applied based on validation performance to mitigate overfitting.

A6. Sampling window, inference aggregation, and thresholds

Each training sample consists of a local temporal context of 10 consecutive profiles (5-min resolution), which is sufficient to
670 exploit short-term continuity without requiring long uninterrupted sequences. During training we use a non-overlapping stride of 10 to avoid excessive redundancy. During inference, we use a stride of 1 and aggregate overlapping predictions by averaging to obtain smoother and more temporally consistent cloud-layer masks.

A probability threshold of 0.5 is used during training. For final evaluation after Stage 3, a slightly more conservative threshold of 0.65 is used to reduce false positives in operational conditions.

A7. Evaluation metrics and dataset split

Cloud-layer identification is evaluated using pixel-wise intersection-over-union (IoU), F1 score, precision, recall, and accuracy. Let TP, FP, TN, and FN denote the number of true positives, false positives, true negatives, and false negatives, respectively, computed from the binary cloud-layer mask over the evaluation set. The metrics are defined as:

$$\text{Precision} = \frac{\text{TP}}{\text{TP} + \text{FP}}, \quad (3)$$

680 $\text{Recall} = \frac{\text{TP}}{\text{TP} + \text{FN}}, \quad (4)$

$$\text{F1} = \frac{2 \text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}} = \frac{2\text{TP}}{2\text{TP} + \text{FP} + \text{FN}}, \quad (5)$$

$$\text{IoU} = \frac{\text{TP}}{\text{TP} + \text{FP} + \text{FN}}, \quad (6)$$

$$\text{Accuracy} = \frac{\text{TP} + \text{TN}}{\text{TP} + \text{TN} + \text{FP} + \text{FN}}, \quad (7)$$

Unless otherwise stated, TP/FP/TN/FN are accumulated over all pixels of the entire test set (i.e., micro-averaging). The
685 dataset is split into training (approximately 80%), validation (approximately 20%), and an independent test set composed of



the three held-out sites (51628, 53845, 58847). The held-out sites are excluded from both pretraining and supervised training to ensure a strict cross-site and cross-time evaluation, as described in the main text.

A8. Special constraints and practical considerations

Two practical constraints are explicitly handled. First, the near-range blind zone is discouraged during training (Stage 2) to prevent non-physical detections. Second, in scenes dominated by aerosols, the distillation objective introduces an aerosol suppression term based on low teacher cloud-score regions to reduce false positives. For rare all-negative cases (no cloud pixels in the reference mask), metrics are computed with a consistent convention to avoid undefined precision/recall.

Code and data availability

The training and inference code, pretrained model weights, and transferable application examples will be made publicly available upon publication. The data and scripts required to reproduce the main figures and tables will be made available to reviewers during the review process upon reasonable request. The full CARLNET observational data used in this study are subject to institutional data-management policies. These data will be released in subsequent data products together with the planned cloud-statistics and optical-parameter-retrieval datasets, where permitted by the relevant data-sharing policy.

Author contributions

WZ and ZY conceived the study and designed the research framework. WZ developed the cloud-layer identification model, implemented the training and inference workflow, performed the experiments, analysed the results, and prepared the initial manuscript draft. ZY supervised the study, contributed to the lidar-processing methodology and experimental design, and revised the manuscript. TS contributed to data processing, model implementation, and experimental analysis. ZB and YD contributed to CARLNET data resources, instrument-operation information, and interpretation of operational lidar observations. HW contributed to cloud and aerosol remote-sensing interpretation and validation discussion. LW contributed to lidar-methodology discussion and interpretation of Raman-polarization lidar observations. YH contributed to radiosonde-related analysis and cloud-structure interpretation. ShL and JC contributed to data organization, annotation support, and experimental checking. XW and SiL contributed to project supervision, resources, and scientific discussion. DM contributed to conceptualization, lidar-remote-sensing interpretation, supervision, and manuscript revision. All authors reviewed and approved the final manuscript.

Competing interests

The authors declare that they have no conflict of interest.



Disclaimer

715 Copernicus Publications remains neutral with regard to jurisdictional claims made in the text, published maps, institutional affiliations, or any other geographical representation in this paper. While Copernicus Publications makes every effort to include appropriate place names, the final responsibility lies with the authors. Views expressed in the text are those of the authors and do not necessarily reflect the views of the publisher.

Acknowledgements

720 The authors thank the China Meteorological Administration and the CARLNET operation teams for maintaining the ground-based Raman lidar observations used in this study. The authors also acknowledge the support from the School of Remote Sensing and Information Engineering, Wuhan University, and the collaborating institutions involved in lidar operation, data processing, and scientific discussion. The radiosonde observations used for independent consistency checks are also gratefully acknowledged.

Financial support

725 This work was supported by the National Key Research and Development Program of China (grant no. 2023YFC3007802), the National Natural Science Foundation of China (grant nos. 42575141, 42475147, 42575138, and 62275202), and the Fundamental Research Funds for the Central Universities (grant no. 20251504004).

Review statement

730 The review statement will be added by Copernicus Publications listing the handling editor as well as all contributing referees according to their status anonymous or identified.

References

- Alizadeh-Choobari, O. and Gharaylou, M.: Aerosol impacts on radiative and microphysical properties of clouds and precipitation formation, *Atmos. Res.*, 185, 53–64, <https://doi.org/10.1016/j.atmosres.2016.10.021>, 2017.
- Ansmann, A., Riebesell, M., and Weitkamp, C.: Measurement of atmospheric aerosol extinction profiles with a Raman lidar, *Opt. Lett.*, 15, 746–748, <https://doi.org/10.1364/OL.15.000746>, 1990.
- 735 Baars, H., et al.: An overview of the first decade of PollyNET: An emerging network of automated Raman-polarization lidars for continuous aerosol profiling, *Atmos. Chem. Phys.*, 16, 5111–5137, <https://doi.org/10.5194/acp-16-5111-2016>, 2016.



- Baars, H., et al.: Target categorization of aerosol and clouds by continuous multiwavelength-polarization lidar measurements, *Atmos. Meas. Tech.*, 10, 3175–3201, <https://doi.org/10.5194/amt-10-3175-2017>, 2017.
- Binietoglou, I., D’Amico, G., Baars, H., Belegante, L., and Marinou, E.: A methodology for cloud masking uncalibrated lidar signals, *EPJ Web Conf.*, 176, 05048, <https://doi.org/10.1051/epjconf/201817605048>, 2018.
- Bu, Z., Dai, Y., Mao, S., Wang, Q., Yin, Z., Wei, Y., Wang, X., Chen, Y., and Zhang, P.: China Aerosol Raman Lidar Network (CARLNET)—Part II: Quality Assessment of Lidar Raw Data, *Remote Sens.*, 18, 663, <https://doi.org/10.3390/rs18040663>, 2026.
- Campbell, J. R., et al.: Full-time, eye-safe cloud and aerosol lidar observation at Atmospheric Radiation Measurement Program sites: Instruments and data processing, *J. Atmos. Ocean. Technol.*, 19, 431–442, [https://doi.org/10.1175/1520-0426\(2002\)019<0431:FTESCA>2.0.CO;2](https://doi.org/10.1175/1520-0426(2002)019<0431:FTESCA>2.0.CO;2), 2002.
- Chisari, A. B., et al.: On the cloud detection from backscattered images generated from a lidar-based ceilometer: Current state and opportunities, in: *Proc. IEEE Int. Conf. Image Process. (ICIP)*, 144–150, <https://doi.org/10.1109/ICIP51287.2024.10647352>, 2024.
- Chisari, A. B., et al.: Benchmarking computer vision architectures for cloud detection from lidar ceilometer backscatter data, *Vis. Comput.*, 41, 9441–9458, <https://doi.org/10.1007/s00371-025-03960-3>, 2025.
- Cromwell, E. and Flynn, D.: Lidar cloud detection with fully convolutional networks, in: *Proc. IEEE Winter Conf. Appl. Comput. Vis. (WACV)*, 619–627, <https://doi.org/10.1109/WACV.2019.00071>, 2019.
- Fernald, F. G.: Analysis of atmospheric lidar observations: Some comments, *Appl. Opt.*, 23, 652–653, <https://doi.org/10.1364/AO.23.000652>, 1984.
- Floutsi, A. A., et al.: DeLiAn – a growing collection of depolarization ratio, lidar ratio and Ångström exponent for different aerosol types and mixtures from ground-based lidar observations, *Atmos. Meas. Tech.*, 16, 2353–2379, <https://doi.org/10.5194/amt-16-2353-2023>, 2023.
- Fu, C., Yang, Z., and Ji, C.: Aerosol identification based on Attention-UNet neural network, *Chin. J. Lasers*, 52, 1010001, <https://doi.org/10.3788/CJL241202>, 2025.
- He, K., et al.: Masked autoencoders are scalable vision learners, *arXiv [preprint]*, <https://doi.org/10.48550/arXiv.2111.06377>, 2021.
- He, Y., Jing, D., Yin, Z., Ohneiser, K., and Yi, F.: Long-term (2010–2021) lidar observations of stratospheric aerosols in Wuhan, China, *Atmos. Chem. Phys.*, 24, 11431–11450, <https://doi.org/10.5194/acp-24-11431-2024>, 2024.
- Illingworth, A. J., et al.: Continuous evaluation of cloud profiles in seven operational models using ground-based observations, *Bull. Amer. Meteor. Soc.*, 88, 883–898, <https://doi.org/10.1175/BAMS-88-6-883>, 2007.
- Jiang, Y., Yao, H., Tao, S., and Liang, J.: Gated skip-connection network with adaptive upsampling for retinal vessel segmentation, *Sensors*, 21, 6177, <https://doi.org/10.3390/s21186177>, 2021.
- Jing, D., et al.: Evolution of tropospheric aerosols over central China during 2010–2024 as observed by lidar, *Atmos. Chem. Phys.*, 25, 17047–17067, <https://doi.org/10.5194/acp-25-17047-2025>, 2025.



- Kalesse-Los, H., et al.: Evaluating cloud liquid detection against Cloudnet using cloud radar Doppler spectra and neural networks, *Atmos. Meas. Tech.*, 15, 279–295, <https://doi.org/10.5194/amt-15-279-2022>, 2022.
- 775 Klett, J. D.: Stable analytical inversion solution for processing lidar returns, *Appl. Opt.*, 20, 211–220, <https://doi.org/10.1364/AO.20.000211>, 1981.
- Loshchilov, I. and Hutter, F.: Decoupled weight decay regularization, in: *Proc. Int. Conf. Learn. Represent. (ICLR)*, <https://doi.org/10.48550/arXiv.1711.05101>, 2019.
- Mona, L., et al.: Lidar measurements for desert dust characterization: An overview, *Adv. Meteorol.*, 2012, 356265, <https://doi.org/10.1155/2012/356265>, 2012.
- 780 Mylonaki, M., et al.: Aerosol type classification analysis using EARLINET multiwavelength and depolarization lidar observations, *Atmos. Chem. Phys.*, 21, 2211–2227, <https://doi.org/10.5194/acp-21-2211-2021>, 2021.
- Oktay, O., et al.: Attention U-Net: Learning where to look for the pancreas, *arXiv [preprint]*, <https://doi.org/10.48550/arXiv.1804.03999>, 2018.
- 785 Pal, S. R., Steinbrecht, W., and Carswell, A. I.: Automated method for lidar determination of cloud-base height and vertical extent, *Appl. Opt.*, 31, 1488–1494, <https://doi.org/10.1364/AO.31.001488>, 1992.
- Ronneberger, O., Fischer, P., and Brox, T.: U-Net: Convolutional networks for biomedical image segmentation, *arXiv [preprint]*, <https://doi.org/10.48550/arXiv.1505.04597>, 2015.
- Sassen, K.: The polarization lidar technique for cloud research: A review and current assessment, *Bull. Amer. Meteor. Soc.*, 72, 1848–1866, [https://doi.org/10.1175/1520-0477\(1991\)072<1848:TPLTFC>2.0.CO;2](https://doi.org/10.1175/1520-0477(1991)072<1848:TPLTFC>2.0.CO;2), 1991.
- 790 Shao, N., et al.: China Aerosol Raman Lidar Network (CARLNET)—Part I: Water vapor Raman channel calibration and quality control, *Remote Sens.*, 17, 414, <https://doi.org/10.3390/rs17030414>, 2025.
- Thorsen, T. J., et al.: Automated retrieval of cloud and aerosol properties from the ARM Raman lidar. Part I: Feature detection, *J. Atmos. Ocean. Technol.*, 32, 1977–1998, <https://doi.org/10.1175/JTECH-D-14-00150.1>, 2015.
- 795 Vainiger, A., et al.: Supervised learning calibration of an atmospheric lidar, in: *Proc. IEEE Int. Geosci. Remote Sens. Symp. (IGARSS)*, 7317–7320, <https://doi.org/10.1109/IGARSS46834.2022.9883078>, 2022.
- Vincent, P., et al.: Extracting and composing robust features with denoising autoencoders, in: *Proc. Int. Conf. Mach. Learn. (ICML)*, 1096–1103, <https://doi.org/10.1145/1390156.1390294>, 2008.
- Voudouri, K. A., et al.: Variability in cirrus cloud properties using a PollyXT Raman lidar, *Atmos. Chem. Phys.*, 20, 4427–4444, <https://doi.org/10.5194/acp-20-4427-2020>, 2020.
- 800 Wang, Y., et al.: Improved retrieval of cloud base heights from ceilometer using a non-standard instrument method, *Atmos. Res.*, 202, 148–155, <https://doi.org/10.1016/j.atmosres.2017.11.021>, 2018.
- Wang, Z. and Sassen, K.: Cloud type and macrophysical property retrieval using multiple remote sensors, *J. Appl. Meteorol. Climatol.*, 40, 1665–1682, [https://doi.org/10.1175/1520-0450\(2001\)040<1665:CTAMPR>2.0.CO;2](https://doi.org/10.1175/1520-0450(2001)040<1665:CTAMPR>2.0.CO;2), 2001.
- 805 Wijnands, J. S., et al.: Deep-Pathfinder: A boundary layer height detection algorithm based on image segmentation, *Atmos. Meas. Tech.*, 17, 3029–3045, <https://doi.org/10.5194/amt-17-3029-2024>, 2024.



- Xie, Q., et al.: Self-training with noisy student improves ImageNet classification, in: Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit. (CVPR), 10684–10695, <https://doi.org/10.1109/CVPR42600.2020.01070>, 2020.
- Yin, Z., Baars, H., Seifert, P., and Engelmann, R.: Automatic lidar calibration and processing program for multiwavelength Raman polarization lidar, EPJ Web Conf., 237, 08007, <https://doi.org/10.1051/epjconf/202023708007>, 2020.
- Zhang, J., et al.: Cloud properties under different synoptic circulations: Comparison of radiosonde and ground-based active remote sensing measurements, Atmosphere, 7, 154, <https://doi.org/10.3390/atmos7120154>, 2016.
- Zhao, C., et al.: Toward understanding of differences in current cloud retrievals of ARM ground-based measurements, J. Geophys. Res. Atmos., 117, D10206, <https://doi.org/10.1029/2011JD016792>, 2012.
- 815 Zhao, C., et al.: A new cloud and aerosol layer detection method based on micropulse lidar measurements, J. Geophys. Res. Atmos., 119, 6788–6802, <https://doi.org/10.1002/2014JD021760>, 2014.
- Zou, W., et al.: Robust lidar–radar composite cloud boundary detection method with rainfall pixels removal, IEEE Trans. Geosci. Remote Sens., 62, 1–16, <https://doi.org/10.1109/TGRS.2024.3476127>, 2024.