

Team GPLSI at QCLEF 2026: Quantum-Inspired Instance Selection and Clustering

Notebook for the QuantumCLEF Lab at CLEF 2026

Santiago Galiano-Segura^{1,4*,†}, Francisco Valero-Abellón^{1,4,†}, Mónica Sánchez-Bellido^{1,†}, Jairo Madrigal-Cutillas^{1,†}, Juan Pablo Consuegra-Ayala^{1,3,4}, Olga Francés^{1,3,4}, Manuel Palomar^{1,3,4}, Paloma Moreda^{1,3,4} and Elena Lloret^{1,3,4}

¹Natural Language Processing and Information Systems Group (GPLSI), University of Alicante, Alicante, Spain, 03690.

²Digital Intelligence Centre (CENID), University of Alicante, Alicante, Spain, 03690.

³University Institute for Computing Research (IUII), University of Alicante, Alicante, Spain, 03690.

⁴Department of Software and Computing Systems (DLSI), University of Alicante, Alicante, Spain, 03690.

Abstract

This paper presents the participation of the GPLSI team in the *Quantum CLEF (QClef) Lab* at *CLEF 2026*, focusing on *Task 2 - Instance Selection* and *Task 3 - Clustering*. The QClef Lab explores the applicability of quantum and quantum-inspired techniques to core AI tasks, emphasizing optimization efficiency and data reduction. For Task 2, we propose two multi-paradigm approaches for selecting representative training instances for sentiment classification: a precision-driven, SVM-augmented LocalSets method, and a compression-driven Reinforcement Learning approach based on stratified topological partitioning. For Task 3, we introduce a unified quantum-inspired clustering framework that integrates four distinct pivot selection strategies for document grouping in embedding space. Our methods achieved competitive performance across both tasks, highlighting a fundamental trade-off between predictive accuracy and computational efficiency. Notably, our *LocalSets* approach achieved the second highest classification effectiveness for instance selection, while our novel *Query-aware Multilevel Graph Clustering* strategy delivered the most robust results for Task 3, effectively balancing ranking preservation (nDCG@10) and cluster compactness (DBI). Overall, our findings support the potential of annealing-based techniques to deliver effective trade-offs between performance and computational efficiency in realistic machine learning pipelines.

Keywords

Quantum Computing, Quantum Annealing, Quantum NLP, CLEF, Instance Selection, Clustering

1. Introduction

The *Quantum CLEF (QClef) Lab* at *CLEF 2026* [1, 2] explores the integration of quantum computing technologies into real-world information retrieval and machine learning workflows. The lab serves as a benchmark initiative to assess how quantum-inspired and quantum-native methods can enhance computational efficiency and effectiveness across core AI tasks. By leveraging both quantum and simulated annealers, QClef aims to address the suitability of these technologies for structured and unstructured data analysis.

QClef 2026 is organized into three primary tasks: *Task 1 - Feature Selection*, *Task 2 - Instance Selection*, and *Task 3 - Clustering*. Each task presents a different machine learning challenge, offering opportunities

CLEF 2026 Working Notes, 21 – 24 September 2026, Jena, Germany

*Corresponding author.

†These authors contributed equally.

✉ santiago.galiano@ua.es (S. Galiano-Segura); francisco.valero@ua.es (F. Valero-Abellón); msb90@gcloud.ua.es (M. Sánchez-Bellido); jmcm39@gcloud.ua.es (J. Madrigal-Cutillas); juan.consuegra@ua.es (J.P. Consuegra-Ayala); olga.frances@ua.es (O. Francés); mpalomar@dlsi.ua.es (M. Palomar); moreda@dlsi.ua.es (P. Moreda); elloret@dlsi.ua.es (E. Lloret)

ORCID: 0009-0008-4291-6152 (S. Galiano-Segura); 0009-0002-4739-2171 (F. Valero-Abellón); 0009-0003-6307-2204 (M. Sánchez-Bellido); 0009-0002-1694-8288 (J. Madrigal-Cutillas); 0000-0003-2009-393X (J.P. Consuegra-Ayala); 0000-0001-6548-5684 (O. Francés); 0000-0002-1441-7865 (M. Palomar); 0000-0002-7193-1561 (P. Moreda); 0000-0002-2926-294X (E. Lloret)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

to apply quantum approaches in data preprocessing, optimization, and representation learning.

Our team, *GPLSI*, following last edition experience [3], participated in *Task 2* and *Task 3*. The main hypothesis guiding our work is that quantum-inspired optimization methods can lead to more representative and compact data subsets, thereby improving downstream learning and analysis tasks such as classification and clustering. In particular, we explore the utility of Quadratic Unconstrained Binary Optimization (QUBO) formulations and hybrid clustering strategies to address instance selection and structure discovery in sentiment analysis datasets.

Our specific objectives are:

- To develop and compare two instance selection approaches for Task 2.
- To design a quantum-inspired clustering method for Task 3.
- To evaluate all methods using the official datasets and protocols provided by QCLEF 2026.

The remainder of this report is structured as follows.

Sections 2 and 3 detail our methodology for instance selection and clustering, respectively. Section 4 presents the experimental setup and discusses results and observations. Finally, Section 5 concludes with a summary and future directions.

2. Proposed approaches for Task 2: Instance Selection

Task 2 focuses on selecting a representative subset of Vader NYT[4] document embeddings to reduce the computational costs of fine-tuning Large Language Models (LLMs) while preserving classification effectiveness. Given N document vectors $x_i \in \mathbb{R}^d$ and a corresponding set of noisy labels, we develop two different approaches in order to use quantum annealing to identify a high-quality, reduced subset that filters out noise and maintains strong predictive power.

2.1. Approach I — SVM-Augmented LocalSets

This methodology builds on the LocalSets instance selection framework derived from one of our *QCLEF 2025* [5] challenge approaches described in [3], extending it with a *SupportVectorMachine(SVM)-based margin recovery* stage and a second *QUBO refinement* pass to recover boundary-critical instances that the geometric pipeline may discard.

2.1.1. Starting Point: LocalSets + QUBO

The starting point [3] is a five-stage hybrid method that combines local set-based instance selection heuristics [6, 7] with quantum annealing. It relies on a taxonomy of instances defined by their role in classification: *noise instances* disagree with their neighborhoods and are removed; *border instances* lie near class boundaries and are unconditionally retained; and *central instances* are prototypical examples of a local region and are subject to QUBO-based selection. The pipeline operates as follows.

1. **Noise removal (LSSm).** The *Local Set-Based Smoother* [6] discards instances whose *harmfulness*—number of times they act as the nearest enemy¹ of another instance—exceeds their *usefulness*—number of local sets they belong to.
2. **Border detection (LSBo).** The *Local Set Border Selector* [6] designates, within each local set, the instance with the lowest *Local Set Cardinality (LSC)* as a border instance. All borders are unconditionally added to the submission set S , and all non-bordered elements are tagged as such.
3. **Clustering (Modified LS-Clustering).** Non-border, non-noisy instances are grouped using a modified local set-based clustering [7] that processes local sets in ascending LSC order, producing compact same-class patches.

¹In this context, an “enemy” refers to an instance belonging to a different class.

4. **Cluster preselection.** For each cluster c_i with centroid μ_i , the cosine distance to its nearest same-class border $b_i^* \in B$ is computed: $d_i^* = \min_{b \in B, y_b = y_i} \left(1 - \frac{\mu_i \cdot b}{\|\mu_i\| \|b\|}\right)$. The $P \leq 150$ clusters with the smallest d_i^* values are retained as QUBO decision variables.
5. **Quantum selection.** The separation margin between preselected clusters is $\delta_{ij} = \|\mu_i - \mu_j\| - (d_i^* + d_j^* + T)$, with tolerance $T = 0.2$. The QUBO matrix $Q \in \mathbb{R}^{P \times P}$ penalises redundant same-class pairs ($\delta_{ij} < 0$) and rewards well-separated ones ($\delta_{ij} \geq 0$):

$$Q[i, j] = \begin{cases} \text{self_bias} & i = j, \\ \alpha \cdot \left(1 - \frac{\mu_i \cdot \mu_j}{\|\mu_i\| \|\mu_j\|}\right)^{-1} & \text{if } \delta_{ij} < 0 \text{ and } y_i = y_j, \\ -\beta \cdot \left(1 - \frac{\mu_i \cdot \mu_j}{\|\mu_i\| \|\mu_j\|}\right)^{-1} & \text{if } \delta_{ij} \geq 0 \text{ and } y_i = y_j, \\ 0 & y_i \neq y_j \end{cases}$$

with $\alpha = \beta = 1.0$, $\text{self_bias} = -1.0$. A class-balance penalty $Q \leftarrow Q + \gamma \mathbf{w} \mathbf{w}^\top$ ($w_i = 1 - 2y_i$, $\gamma = 5.0$) promotes equal class representation. Simulated Annealing(SA) or Quantum Annealing(QA) then selects the optimal cluster subset, and members of selected clusters are added to S .

2.1.2. Stage I: SVM Margin Recovery

A known limitation of geometry-based local set criteria is that instances near the decision boundary may be discarded as locally redundant despite being critical for margin definition. To recover these instances, we train a SVM with RBF kernel ($C = 1.0$, $\gamma = \text{scale}$) on the initial selection S . Let $f(x)$ denote the SVM decision function. The margin threshold τ is estimated as the mean absolute decision score over the free support vectors (those with $0 < \alpha_i < C$):

$$\tau = \frac{1}{|SV_{\text{free}}|} \sum_{i \in SV_{\text{free}}} |f(x_i)|.$$

Each discarded instance $x \in D = \mathcal{X} \setminus S$ is evaluated by $|f(x)|$: instances satisfying $|f(x)| < \tau$ lie closest to the margin and are therefore most informative for boundary definition. The top- K such instances, ranked by ascending $|f(x)|$, are recovered and added to S , yielding the augmented set S' .

2.1.3. Stage II: QUBO Refinement

The recovered instances may introduce redundancy into the selection. To filter out genuinely redundant additions, we rebuild the local set cluster structure on S' : border instances are recomputed via LSBo and the remaining instances are re-clustered following the same Modified LS-Clustering scheme. Since $|S'| \ll |\mathcal{X}|$, the number of resulting clusters is substantially smaller than in the original pass (typically ≤ 150), so all clusters serve directly as QUBO decision variables. The same QUBO formulation is applied to this refined cluster set and solved again with SA or QA. Instances belonging to clusters discarded in this second pass are removed from S' , yielding a compact, margin-aware selection that preserves the boundary-defining instances identified by the SVM while eliminating those that do not meet the geometric QUBO quality criterion.

2.2. Approach II — Stratified Partitioning DQN

This approach addresses the instance selection problem through a methodology that combines topological data partitioning with a reinforcement learning agent for the instance reduction task by defining the problem as Markov Decision Process(MDP) based on the methodology exposed by A. Jha et al.[8]. Our approach adapts this methodology to the high-dimensional document embedding space provided

of the VaderNYT dataset [4] and sentiment classification task, using a selection agent to identify a representative subset of 150 clusters of instances optimized for QUBO formulations.

To provide a clear roadmap of this approach, the execution workflow consists of three primary stages that will be detailed in this section:

1. **Data Partitioning:** Grouping the document embeddings into $K = 227$ discrete clusters using a stratified partitioning strategy.
2. **Agent-Based Selection:** Training a DQN agent to identify the optimal 150 clusters, guided by a reward function conditioned on validation accuracy.
3. **Quantum Optimization:** Applying a QUBO formulation to the 150 RL-selected clusters to extract the final representative subset, framing the spatial diversification of the centroids as a minimum-energy optimization problem.

2.2.1. Markov Decision Process Formulation

To optimize the selection of document embeddings, we formulate the problem as a sequential decision-making process defined by this tuple (S, A, R, γ) .

- **State space (S):** The state $s_k \in \mathbb{R}^{768}$ is defined by the centroid of a candidate cluster partition. This representation captures the local topological features of a specific region in the embedding space.
- **Action space (A):** For each state s_k , the agent performs a binary action $a \in \{0, 1\}$. The action 1 implies that the cluster k is included into the training subset, while 0 implies that the cluster is not taken.
- **Reward function (R):** The reward R_k for selecting a specific cluster partition is defined by its *Proxy utility score*, a supervision-based metric designed to quantify the "informativeness" of a topological data region. While A. Jha, et al. [8] use a LLM as the proxy evaluation function, we adapt the mechanism to the computational constraints and specific objectives of Task 2 by employing a Multi-Layer Perceptron (MLP) Classifier. For each candidate cluster C_k , a MLP model is trained on the documents embeddings with in that partition. The reward R_k is then derived from the model's classification accuracy when evaluated against a validation set.
- **Discount factor (γ):** We set $\gamma = 0.99$ to prioritize long-term representational accuracy across the entire dataset selection.

2.2.2. Data partitioning

To generate the discrete states required for the Markov Decision Process(MDP) and reduce the action space from individual document vectors to representative neighborhoods, we partition the high dimensional embedding space $\mathbb{R}^d$ into K discrete clusters. The number of clusters K is determined dynamically for each fold based on the training set size to ensure $K > 150$. We tested three different clustering strategies to identify the optimal configuration of our RL agent.

Standard K-Means: As fundamental spatial baseline, the framework uses the standard K-Means algorithm to minimize the within-cluster sum of squares (WCSS).

1. Initializes centroids using the KMeans++ scheme to improve convergence stability.
2. Document embeddings are grouped based on their Euclidean distances proximity.

Stratified partitioning: Instead of relying solely on spatial proximity, this sampling methodology adopts a label-centric approach to ensure that the broader dataset's class distribution is faithfully mirrored within every individual subset. By prioritizing this global balance, the strategy structures the partitions so that each selected unit acts as a proportionally accurate, representative sample of the entire dataset.

1. Algorithm iterates over each unique label value $\{0, 1\}$ and identifies all the associated document indices.

2. Random permutation is applied to the indices within each class to ensure that any temporal or source bias in the dataset is mitigated.
3. The indices for each class are then partitioned into K chunks using a balanced split.
4. The chunks are aggregated across all classes to form the final K clusters, ensuring that every partition serves as representative sample of the global label distribution.

Stratified Balanced K-Means: Algorithm that serves as an iterative hybrid approach designed to optimize both spatial coherence and class representation.

1. Generates initial centroids through a standard unsupervised K-Means fit to establish a baseline topology.
2. For each instance, the algorithm calculates the Euclidean distance to all cluster centers. The instance is assigned to the nearest cluster only if the corresponding class quota $q_{k,c}$ has not been reached. If the primary quota is exhausted, the algorithm evaluates the next nearest centroid, ensuring that every instance is assigned to a cluster without violating the balance constraints.
3. The cluster centers are updated by computing the arithmetic mean of all current members:

$$\mu_k = \frac{1}{|C_k|} \sum_{x \in C_k} x.$$
4. The process of assignment and refinement is repeated for a fixed number of iterations ($n_iters = 5$) or until the shift in cluster centers falls below a defined tolerance threshold.

To determine the most effective grouping mechanism and the optimal initial number of clusters, we evaluated each data partitioning strategy with $K \in \{151, 227, 341, 512\}$. These values were selected arbitrarily to assess how different levels of granularity affect model performance: smaller values of K correspond to coarser partitions, whereas larger values produce finer-grained partitions. For each of the five folds provided in the challenge, we extracted a test set and used the same proxy MLP architecture to ensure a fair comparison of generalization capabilities. The following metrics were used to evaluate the agent trained under each partitioning strategy and value of K :

- **Mean test accuracy:** This metric represents the average accuracy achieved by the model on the test set generated from train set provided across all folds. It serves as a direct measure of the generalization capability of the model, with higher values indicating better performance on unseen data.
- **Mean validation accuracy:** This metric represents the average accuracy achieved by the model on the validation set from train set provided across all folds.
- **Mean absolute gap:** The mean absolute gap quantifies the average absolute difference between validation accuracy and test accuracy across all folds. Lower values indicate more stable generalization, as the model’s performance on the validation set closely matches its performance on the test set.

$$\text{MeanAbsoluteGap} = \frac{1}{n} \sum_{i=1}^n |\text{val_acc}_i - \text{test_acc}_i|$$

- **Mean subset ratio:** This metric measures the compression efficiency of the selected clusters relative to the training dataset size. It is defined as:

$$\text{SubsetRatio} = \frac{\text{SubsetSize}}{\text{TrainingSize}}$$

Lower values indicate better compression, as the selected subset is smaller relative to the full training dataset, while still maintaining performance.

- **Mean reward contrast:** Evaluates the difference in average rewards between the selected clusters and the non-selected clusters.

$$\text{RewardContrast} = R_{\text{selected}} - R_{\text{non-selected}}$$

- **Composite score:** This metric is a weighted combination of the above metrics, designed to provide an overall ranking of the methodologies. The weights values were arbitrary defined in order to prioritize accuracy metrics.

$$\text{CompositeScore} = 0.45 \cdot \text{score_test} + 0.30 \cdot \text{score_stability} + 0.15 \cdot \text{score_compression} + 0.10 \cdot \text{score_reward_quality} \quad (1)$$

An analysis of the comparative performance in Table 1 demonstrates that Stratified Partitioning outperforms alternative clustering methods, achieving the highest composite score due to its superior balance of test accuracy and subset compression. Additionally, the state-space optimization results summarized in Table 2 isolate $K = 227$ as the optimal number of clusters. This specific granularity provides a high-signal state representation for the *Deep Q-Network* (DQN) [9] agent, maximizing classification generalization while maintaining a highly controlled validation-to-test gap.

Table 1

Composite Ranking of data partitioning methodologies for approach 2.

Method	Mean Test Acc	Mean Val Acc	Mean Abs Gap	Mean Subset Ratio	Mean Reward Contrast	Composite Score
Stratified Partitioning	0.9427	0.9540	0.0152	0.6012	0.0540	0.6633
Stratified Bal. KM	0.9406	0.9553	0.0165	0.6022	0.1430	0.5629
K-Means	0.9277	0.9342	0.0105	0.7931	0.4140	0.4000

Table 2

Composite ranking of different number of clusters k .

k	Mean Test Acc	Mean Val Acc	Mean Abs Gap	Mean Subset Ratio	Mean Reward Contrast	Composite Score
151	0.9333	0.9404	0.0111	0.9934	0.1260	0.3000
227	0.9475	0.9582	0.0114	0.6608	0.0393	0.8729
341	0.9465	0.9593	0.0162	0.4404	0.0267	0.7698
512	0.9434	0.9582	0.0222	0.3102	0.0241	0.5429

2.2.3. Deep Q-Network Agent Architecture and Policy Optimization

To learn the optimal selection policy over the generated state space, the agent approximates the action-value function $Q(s, a)$ using a neural network and an iterative optimization loop designed to stabilize training under noisy conditions.

- **Neural Function Approximation:** The action-value function is modeled using a MLP regressor configured with two hidden layers of sizes 128 and 64, utilizing Rectified Linear Unit (ReLU) activation functions. The network accepts a 768-dimensional cluster centroid vector as the input state s_k and evaluates the expected cumulative reward for both inclusion ($a = 1$) and exclusion ($a = 0$) actions.
- **Experience Replay Buffer:** To break the temporal correlations inherent in sequential cluster evaluations, transitions $(s, a, r, s', \text{done})$ are collected and stored within a cyclic memory buffer with a capacity of 2,000 episodes. During the optimization step, the agent samples randomized mini-batches of size 32 from the buffer to compute gradient steps, using an initial learning rate of $\alpha = 0.001$.
- **Target Network Synchronization:** To prevent harmful feedback loops and non-stationarity during the bootstrapping phase, we decouple value estimation by maintaining a secondary target

network ($\hat{Q}$). This network remains parameter-locked and is updated with the weights of the active online network every 10 episodes, stabilizing the calculation of the Temporal Difference (TD) target:

$$Y = r + \gamma \max_{a'} \hat{Q}(s', a') \cdot (1 - \text{done})$$

The network minimizes the mean squared Bellman error via partial fitting on the sampled mini-batch.

- **Dynamic Exploration-Exploitation Schedule:** Action selection follows an ϵ -greedy policy where the exploration factor decays from $\epsilon_{\text{start}} = 1.0$ down to a functional floor of $\epsilon_{\text{end}} = 0.05$. The decay parameter is adjusted dynamically relative to the choice of k to guarantee sufficient exploration across the topological partition space:

$$\epsilon_{e+1} = \max \left(\epsilon_{\text{end}}, \epsilon_e \cdot \left(\frac{\epsilon_{\text{end}}}{\epsilon_{\text{start}}} \right)^{\frac{1}{\lceil 0.75 \cdot (100 + \lceil k/2 \rceil \rceil)}} \right)$$

This dynamic horizon guarantees that larger state resolutions receive proportionally longer exploration windows before the policy transitions toward exploitation.

2.2.4. Supervised Cardinality-Constrained Cluster-Level Max-Sum Diversification QUBO

For the pre-filtered candidate partitions consisting of m discrete clusters, we follow a supervised Max-Sum Diversification formulation inspired by the algebraic structures of Bauckhage et al. [10]. We define the Euclidean distance between cluster centroids as:

$$d_{cd} = \|\boldsymbol{\mu}_c - \boldsymbol{\mu}_d\|_2, \quad (2)$$

where $\boldsymbol{\mu}_c \in \mathbb{R}^{768}$ represents the centroid vector of cluster c . The binary decision variable $z_c \in \{0, 1\}$ indicates whether cluster c is selected ($z_c = 1$) or excluded ($z_c = 0$) for the final subset rehydration.

To frame the maximization of class-signal utility and spatial diversification as a minimum-energy problem for the solver, the unconstrained Binary Quadratic Model (BQM) is defined as:

$$E_{\text{base}}(\mathbf{z}) = \sum_{c=1}^m Q_{cc} z_c + \sum_{c < d} Q_{cd} z_c z_d, \quad (3)$$

where the linear terms along the diagonal prioritize partitions governed by frequent global class labels:

$$Q_{cc} = -\frac{N_{\ell_c}}{m}, \quad (4)$$

with N_{ℓ_c} representing the total instance frequency of the dominant label ℓ_c associated with cluster c . The quadratic off-diagonal interaction terms enforce class-conditional geometric constraints:

$$Q_{cd} = \begin{cases} d_{cd}, & \text{if } \ell_c = \ell_d, \\ -d_{cd}, & \text{if } \ell_c \neq \ell_d, \end{cases} \quad (5)$$

where a positive coefficient for identical dominant labels ($\ell_c = \ell_d$) acts as a structural repulsion penalty to suppress informational redundancy within the same sentiment category.

The total constrained objective function incorporating a hard cardinality constraint to select exactly a target number of clusters $k = \lfloor p \cdot m \rfloor$ is expressed as:

$$E(\mathbf{z}) = E_{\text{base}}(\mathbf{z}) + \lambda \left(\sum_{c=1}^m z_c - k \right)^2. \quad (6)$$

3. Proposed approaches for Task 3: Clustering

Task 3 focuses on clustering ANTIQUE document embeddings to improve retrieval efficiency while preserving ranking quality. Given N document vectors $\mathbf{x}_i \in \mathbb{R}^d$ and a target panel size $k \in \{10, 25, 50\}$, we select representative medoids and assign each document to its nearest medoid in cosine space. All methods share the same optimization core and differ only in how the candidate set (or compressed representation) is built before annealing.

3.1. Shared QUBO Core (Bauckhage-Style k -Medoids)

For a candidate set of size $P \leq 150$, we follow a Bauckhage-style binary k -medoids formulation [11]. We define cosine distances

$$\delta_{ij} = 1 - \cos(\hat{\mathbf{x}}_i, \hat{\mathbf{x}}_j), \quad (7)$$

where $\hat{\mathbf{x}}_i$ are ℓ_2 -normalized embeddings. Binary variable $y_i \in \{0, 1\}$ indicates whether candidate i is selected as medoid.

The unconstrained QUBO is

$$E_{\text{base}}(\mathbf{y}) = \sum_{i=1}^P Q_{ii} y_i + \sum_{i < j} Q_{ij} y_i y_j, \quad (8)$$

with

$$Q_{ii} = -\alpha \sum_{j=1}^P \delta_{ij}, \quad Q_{ij} = 2\beta \delta_{ij}, \quad (9)$$

and $\alpha = 1/k$, $\beta = 1/P$. The constrained objective is then:

$$E(\mathbf{y}) = E_{\text{base}}(\mathbf{y}) + \lambda \left(\sum_{i=1}^P y_i - k \right)^2. \quad (10)$$

In implementation, the cardinality term is instantiated with `dimod.generators.combinations`, and λ is set adaptively from the unconstrained BQM energy scale (`maximum_energy_delta`).

3.2. Unified Execution Pipeline

Each run follows the same six-stage workflow, regardless of the candidate generation strategy:

Algorithm 1: Unified Task 3 Pipeline **Input:** document embeddings X , target k , method-specific configuration. **Output:** k clusters and k final medoids.

1. Build representative set R (or compressed graph/spectral representation), with $|R| \leq 150$.
2. Build the QUBO/BQM on R with exact- k cardinality constraint.
3. Solve with SA (or QA when enabled) and extract selected indices.
4. Map selected representatives to corpus medoid candidates.
5. Assign all documents to nearest medoid in cosine space and enforce exactly k clusters.
6. Recompute final medoids from cluster assignments and export the official submission format.

3.3. Approach 1: Farthest-Point Sampling

Farthest-Point Sampling (FPS) is our unsupervised geometric baseline. It greedily builds a pivot set by repeatedly selecting the point farthest (in max-min sense) from already selected pivots [12]. Given pivot budget $P \in \{100, 150\}$, we first choose an extreme point relative to the corpus centroid, then iteratively maximize diversity:

$$p_{t+1} = \arg \max_i \left(1 - \max_{p \in S_t} \cos(\hat{\mathbf{x}}_i, \hat{\mathbf{x}}_p) \right). \quad (11)$$

The selected pivots feed the shared QUBO core. Its dominant cost is $\mathcal{O}(NP)$.

3.4. Approach 2: Query-Aligned FPS

Query-Aligned FPS is an extension proposed in this work: it keeps the FPS geometric core and adds a query-aware relevance–diversity trade-off, conceptually related to MMR-style reranking [13]. For each document, we compute:

$$q_i = \max_{\mathbf{q} \in \mathcal{Q}} \cos(\hat{\mathbf{x}}_i, \hat{\mathbf{q}}), \quad (12)$$

then normalize q_i to $[0, 1]$. At each iteration, selection maximizes:

$$\text{score}(i) = (1 - \omega) \cdot \text{diversity}(i) + \omega \cdot q_i, \quad (13)$$

with $\omega = 0.5$ in the selected configurations and $P \in \{125, 150\}$. This preserves geometric coverage while biasing pivots toward query-relevant regions. Its dominant cost is $\mathcal{O}(NP + N|\mathcal{Q}|)$.

3.5. Approach 3: Query-aware Multilevel Graph Clustering

Query-aware Multilevel Graph Clustering (QMGC) is also proposed in this work and combines a query-aware KNN graph with multilevel graph reduction ideas [14] before annealing:

1. Build R representative points via MiniBatchKMeans [15] (up to 150 representatives).
2. Compute query alignment scores q_i on representatives.
3. Build weighted KNN graph with

$$W_{ij} = (1 - d_{\cos}(i, j)) + \lambda_q \min(q_i, q_j). \quad (14)$$

4. Coarsen the graph greedily by merging strongly connected groups until target size t_c , following the graph-coarsening rationale in [14].
5. Solve a QUBO over coarsened supernodes with centrality/diversity balancing parameters $(\alpha_c, \beta_c, \gamma_c)$.
6. Expand selected supernodes back to original representatives and continue with the shared post-processing pipeline.

This design injects both structural and query-aware signals before optimization. With at most 150 representatives, runtime is dominated by KNN graph construction and pairwise group comparisons during coarsening.

3.6. Approach 4: Quantum Clustering via Spectral Seeding

Quantum Clustering via Spectral Seeding (QCSS) is also proposed in this work and combines geometric coverage with graph spectral structure [16, 17, 18]:

1. Select up to 150 representatives with FPS.
2. Build cosine KNN graph on representatives.
3. Compute normalized graph Laplacian

$$L = I - D^{-1/2} W D^{-1/2}. \quad (15)$$

4. Extract the first d_{spec} eigenvectors (with $d_{\text{spec}} \approx 1.5k$) to obtain spectral embedding Z .
5. Solve the shared Bauckhage QUBO on Z instead of raw embeddings.

Operating in spectral space helps separate community structure prior to medoid selection. Its dominant step is eigendecomposition on the representative graph (approximately cubic in representative-set size).

3.7. Implementation Notes

All runs use the same targets $k \in \{10, 25, 50\}$ and fixed random seed for reproducibility. SA is always available and is executed with a fixed read budget; QA follows the same pipeline when QPU execution is enabled. After annealing, we enforce exactly k clusters, recompute final medoids from assigned documents, and export submissions in the official Task 3 format. The next section reports the empirical comparison across methods and k values using the official evaluation protocol.

4. Experiments

This section presents the experimental setup and results for our proposed approaches across two main tasks: *Task 2 (Instance Selection)* and *Task 3 (Clustering)*. Based on the official results computed by the challenge organizers, we discuss our observations and evaluate how our methods stand alongside the baselines and competing submissions.

4.1. Results

The evaluation was divided into two main scenarios corresponding to the official tasks defined in the challenge. The selected tasks were: (i) **Task 2: Instance Selection** and (ii) **Task 3: Clustering**. Each task involved the use of SA and QA approaches.

4.1.1. Task 2: Instance Selection

Task 2 focuses on selecting the most representative subsets of training data for fine-tuning a LLM in sentiment classification task, aiming to reduce computational cost while preserving effectiveness. The evaluation considered both the macro F_1 score and the *Average reduction*. The methodologies detailed in Sections 2.1 and 2.2 were executed on both *quantum* and *simulated* annealer. Both approaches were executed with 1000 simulations, with the approach in Section 2.2 additionally executed using 2000 simulations.

Table 3 compares our instance selection approaches: **Local Sets** (Section 2.1, listed as *LocalSets-step-1/2*) and **Stratified Partitioning RL** (Section 2.2, listed as *stratified-partitioning-RL*) with other teams submissions. Results, averaged across five cross-validation folds, are sorted by macro F_1 score to account for class imbalance. *Avg Reduction* indicates the average percentage of retained documents. Furthermore, the table details computational overhead via fine-tuning, prediction, and annealing times, while the *Type* column specifies the optimization paradigm: Simulated Annealing (SA) or Quantum Annealing (QA).

Table 3

Summary of *Task 2 - Instance Selection* results (averaged over 5 folds). Types SA and QA stand for Simulated Annealing and Quantum Annealing

Group	Submission ID	Macro F_1 Avg	Avg Reduction	Avg Fine-Tuning + Prediction Time (s)	Avg Annealing Time (us)	Type
BASLINE	Baseline_without_noise	88.9(0.8)	–	1997.3(5.7)	–	–
BASLINE	Baseline_with_noise	84.3(2.5)	–	1904.3(1.4)	–	–
R2AI	QA_r2ai_001	75.7(3.7)	26.02%	1458.3(235.5)	191481640	H
GPLSI	SA_gplsi_LocalSets-step-1	67.9(10.5)	12.74%	1675.5(330.5)	21914616	SA
GPLSI	SA_gplsi_LocalSets-step-2	67.9(10.5)	12.74%	1675.3(330.2)	25083288	SA
R2AI	SA_r2ai_001	67.8(8.9)	26.02%	1517.7(306.5)	976908763	SA
Entropy Of Void	SA_entropyOfVoid_QuboKMedoids	65.9(10)	59.99%	816.8(0.5)	6392047845	SA
DS@GT QuantumCLEF	QA_ds-at-gt-qclef_svc_cluster_rep	64.8(8.9)	67.62%	678.7(75.2)	88104216	QA
Entropy Of Void	QA_entropyOfVoid_QuboKMedoids	63.9(6.7)	59.83%	819.3(5.4)	21572994	QA
DS@GT QuantumCLEF	SA_ds-at-gt-qclef_svc_cluster_rep	60.4(7.4)	71.44%	609.8(3.4)	(Checking)	SA
FAST-NUCES	QA_FAST-NUCES_DivideMix	58.5(11.9)	60.98%	797.4(11.1)	2643107	QA
FAST-NUCES	SA_FAST-NUCES_Algorithm1	57.9(8.8)	75.03%	541.7(0.3)	167180028	SA
FAST-NUCES	SA_FAST-NUCES_DivideMix	56.3(9.1)	70.04%	642.4(28.6)	2918988448	SA
GPLSI	QA_gplsi_LocalSets-step-1	55.5(8.7)	49.57%	1007(359.6)	997262	QA
GPLSI	QA_gplsi_LocalSets-step-2	55.5(8.7)	49.57%	1060.2(406.2)	1042623	QA
GPLSI	SA_gplsi_stratified-partitioning-RL-simulations-2000	53.3(10.1)	82.42%	409.7(3.9)	182810207	SA
GPLSI	QA_gplsi_stratified-partitioning-RL-simulations-2000	50.4(7.7)	82.30%	412.1(5.4)	2727665	QA
GPLSI	SA_gplsi_stratified-partitioning-RL-simulations-1000	50.2(2.8)	82.43%	409.6(2)	91374301	SA
GPLSI	QA_gplsi_stratified-partitioning-RL-simulations-1000	49.8(9.7)	82.37%	410.7(2.9)	1404424	QA

4.1.2. Task 3: Clustering

Task 3 involves clustering sentence embeddings to support efficient document retrieval and exploration. Quality was evaluated both intrinsically with the Davies-Bouldin Index and extrinsically with $nDCG@10$ on test queries. We submitted clustering results on the large and small variants of the ANTIQUE dataset using both SA and QA. All methods used were executed using 2000 simulations. Table 4 summarizes our submitted runs for Task 3, indicating which approaches were applied to each dataset using SA and QA.

Table 4 reports the evaluation results for Task 3 on the ANTIQUE dataset. It compares the performance of the submitted clustering approaches for different values of k (10, 25, and 50), including the **baseline**, the **GPLSI** submissions, and the best-performing **entropyOfVoid** configurations. Results are evaluated intrinsically using the Davies-Bouldin Index (DBI), where lower values indicate better cluster separation, and extrinsically using $nDCG@10$, which measures the quality of document retrieval in the top 10 results. The table includes executions performed with both Simulated Annealing (SA) and Quantum Annealing (QA), all executed using 2000 simulations. The *Anneal time* column reports the execution time in microseconds for each run, while the *Type* column indicates whether the corresponding submission was executed using SA or QA.

Table 4

Summary of *Task 3 - ANTIQUE dataset* results. Results include both Simulated Annealing (SA) and Quantum Annealing (QA) executions.

k	Group	Submission ID	$nDCG@10 \uparrow$	DBI $\downarrow$	Anneal time (μs)	Type
10	BASELINE	BASELINE_10	0.5509	7.9892	–	–
10	GPLSI	10_QA_gplsi_FPS001	0.4862	7.1169	447596	QA
10	GPLSI	10_SA_gplsi_FPS001	0.4654	7.2543	13703089	SA
10	GPLSI	10_QA_gplsi_QAL003	0.4857	6.9728	490163	QA
10	GPLSI	10_SA_gplsi_QAL003	0.5562	7.8100	19506025	SA
10	GPLSI	10_QA_gplsi_QAL004	0.5003	7.0393	545087	QA
10	GPLSI	10_SA_gplsi_QAL004	0.4956	6.8404	25765877	SA
10	GPLSI	10_QA_gplsi_QMGC005	0.4832	6.8672	404240	QA
10	GPLSI	10_SA_gplsi_QMGC005	0.5628	6.9122	8424728	SA
10	GPLSI	10_QA_gplsi_QMGC006	0.5506	6.6033	407082	QA
10	GPLSI	10_SA_gplsi_QMGC006	0.5564	6.7533	7219938	SA
10	BEST $nDCG@10$	10_QA_entropyOfVoid_G80noPCAeuc	0.6363	6.4572	847635	QA
10	BEST DBI	10_QA_entropyOfVoid_G80noPCAcos	0.5948	6.3292	850233	QA
25	BASELINE	BASELINE_25	0.5284	6.1201	–	–
25	GPLSI	25_QA_gplsi_FPS001	0.4450	5.7445	432802	QA
25	GPLSI	25_SA_gplsi_FPS001	0.4413	5.9002	17131481	SA
25	GPLSI	25_QA_gplsi_QAL003	0.4435	5.8082	490287	QA
25	GPLSI	25_SA_gplsi_QAL003	0.5071	5.6069	24311199	SA
25	GPLSI	25_QA_gplsi_QAL004	0.5054	5.7333	559921	QA
25	GPLSI	25_SA_gplsi_QAL004	0.4861	5.5941	31646655	SA
25	GPLSI	25_QA_gplsi_QMGC005	0.5513	5.2352	326837	QA
25	GPLSI	25_SA_gplsi_QMGC005	0.5799	5.3022	10171741	SA
25	GPLSI	25_QA_gplsi_QMGC006	0.5448	5.1040	344399	QA
25	GPLSI	25_SA_gplsi_QMGC006	0.5165	5.2551	8567676	SA
25	BEST $nDCG@10/DBI$	25_QA_entropyOfVoid_G80noPCAcos	0.5858	5.0280	843446	QA
50	BASELINE	BASELINE_50	0.4656	5.3679	–	–
50	GPLSI	50_QA_gplsi_FPS001	0.4483	4.6322	426281	QA
50	GPLSI	50_SA_gplsi_FPS001	0.4608	4.8406	16852807	SA
50	GPLSI	50_QA_gplsi_QAL003	0.4522	4.9158	476045	QA
50	GPLSI	50_SA_gplsi_QAL003	0.4540	4.9734	29252806	SA
50	GPLSI	50_QA_gplsi_QAL004	0.4377	4.9023	557730	QA
50	GPLSI	50_SA_gplsi_QAL004	0.4073	4.9647	39436923	SA
50	GPLSI (BEST DBI)	50_QA_gplsi_QMGC005	0.5547	4.2726	454761	QA
50	GPLSI	50_SA_gplsi_QMGC005	0.5360	4.4472	11092752	SA
50	GPLSI (BEST $nDCG@10$)	50_QA_gplsi_QMGC006	0.5593	4.3579	450156	QA
50	GPLSI	50_SA_gplsi_QMGC006	0.5509	4.4169	12285791	SA

4.2. Discussion

This section discusses the main findings of our participation in the QuantumCLEF Lab, organized by task. In Sections 4.2.1 and 4.2.2, we discuss the outcomes of the instance selection and clustering tasks. Our analysis emphasizes the effectiveness and operational trade-offs of our approaches, benchmarking them against both the baseline and other participating submissions.

4.2.1. Task 2: Instance Selection

As shown in Table 3, we observe a clear correlation between the *Avg Reduction* and *Macro F_1* metrics. Specifically, more aggressive instance selection generally leads to a decrease in model performance.

Our **Local Sets** approach prioritized classification effectiveness, obtaining the highest macro F_1 score among our submissions (67.9) and outperforming several competing methodologies. However, its overall data reduction rate was comparatively conservative at 12.74%. Furthermore, the identical reduction rates observed after both the geometric phase (step-1) and the SVM-augmented phase (step-2) indicate that the two-step pipeline did not work as intended to further compress the dataset. Additionally, a striking difference in performance emerged between the Simulated Annealing (SA) and Quantum Annealing (QA) executions for this method. The QA solver applied a much more aggressive data selection strategy than its SA counterpart, which consequently degraded the model’s downstream classification performance.

On the other hand, for the **Stratified Partitioning RL** methodology, we observe an aggressive data compression strategy that yields the highest average reduction rates across all submissions, consistently exceeding 82%. This aggressive data dropout can be primarily attributed to the granularity of the state space in which the DQN agent was trained—specifically, the number of discrete target clusters (k). By forcing the agent to evaluate and select from a highly compressed topological representation, the maximum proportion of retained instances is inherently bottlenecked by this initial partitioning phase. While this substantial data pruning predictably results in lower macro F_1 scores (ranging from 49.8 to 53.3) compared to the **Local Sets** approach and the baselines, it successfully addresses the core objective of minimizing computational overhead. Notably, the *Stratified Partitioning RL* runs achieved the lowest combined fine-tuning and prediction times in the entire evaluation (approximately 410 seconds). This represents a nearly 80% reduction in downstream processing time compared to the noise-free baseline, highlighting a stark but functional trade-off between predictive accuracy and operational efficiency. Furthermore, scaling the execution from 1000 to 2000 simulations resulted in measurable improvements in the macro F_1 score, particularly for the SA execution (increasing from 50.2 to 53.3). This trend suggests that allowing the solver a larger sampling budget effectively enhances its ability to explore the complex, high-dimensional partitions formulated by the agent.

To complement the metrics reported in the challenge, we propose a metric that measures the performance efficiency per retained data instance, benchmarking it against the baseline with noise injected samples that uses the entire dataset.

$$\text{Data Efficiency} = \frac{\text{Macro } F_1}{100 - \text{Avg Reduction}}$$

Essentially, this formulation quantifies the predictive density of the selected subset. A lower score indicates low data efficiency, meaning a larger volume of data was required to achieve the resulting F_1 score. Conversely, a higher score denotes high efficiency, demonstrating that the model achieved its classification performance using only a minimal data footprint.

In Table 5, we analyze both the absolute and relative data efficiencies across all submissions. By normalizing the efficiency against the official baseline, the relative metric directly quantifies how much each method improves the predictive yield per retained instance. This comparison clearly demonstrates that our *Stratified Partitioning RL* configurations achieve the highest efficiency ratios overall—reaching up to 3.60 times the baseline efficiency. Ultimately, this confirms that while the method understandably ranks lower in raw macro F_1 scores, it successfully maximizes data compression and drastically minimizes the downstream computational footprint.

Table 5

Summary of *Task 2 - Instance Selection* results, highlighting absolute and relative Data Efficiency alongside Macro F_1 and Average Reduction rankings.

Group	Submission id	Data Efficiency (abs)	Data Efficiency (rel)	Ranking Macro F_1 Avg	Ranking Average Reduction	Type
GPLSI	SA_gplsi_stratified-partitioning-RL-simulations-2000	3.03	3.60	15	2	SA
GPLSI	SA_gplsi_stratified-partitioning-RL-simulations-1000	2.86	3.39	17	1	SA
GPLSI	QA_gplsi_stratified-partitioning-RL-simulations-2000	2.85	3.38	16	4	QA
GPLSI	QA_gplsi_stratified-partitioning-RL-simulations-1000	2.82	3.35	18	3	QA
FAST-NUCES	SA_FAST-NUCES_Algorithm1	2.32	2.75	11	5	SA
DS@GT QuantumCLEF	SA_ds-at-gt-qclef_svc_cluster_rep	2.11	2.51	9	6	SA
DS@GT QuantumCLEF	QA_ds-at-gt-qclef_svc_cluster_rep	2.00	2.37	7	8	QA
FAST-NUCES	SA_FAST-NUCES_DivideMix	1.88	2.23	12	7	SA
Entropy Of Void	SA_entropyOfVoid_QuboKMedoids	1.65	1.95	6	10	SA
Entropy Of Void	QA_entropyOfVoid_QuboKMedoids	1.59	1.89	8	11	QA
FAST-NUCES	QA_FAST-NUCES_DivideMix	1.50	1.78	10	9	QA
GPLSI	QA_gplsi_LocalSets-step-1	1.10	1.31	13	12	QA
GPLSI	QA_gplsi_LocalSets-step-2	1.10	1.31	13	12	QA
R2AI	QA_r2ai_001	1.02	1.21	2	14	H
R2AI	SA_r2ai_001	0.92	1.09	5	14	SA
GPLSI	SA_gplsi_LocalSets-step-1	0.78	0.92	3	16	SA
GPLSI	SA_gplsi_LocalSets-step-2	0.78	0.92	3	16	SA
BASELINE	Baseline_with_noise	0.84	1.00	1	18	-

Overall, the results from Task 2 highlight a fundamental trade-off between predictive accuracy and computational efficiency in LLM fine-tuning. Our two proposed methodologies successfully navigate opposite ends of this spectrum: the *Local Sets* approach acts as a precision tool, maintaining high classification fidelity with moderate data pruning, whereas the *Stratified Partitioning RL* serves as an aggressive compression method, trading a portion of effectiveness for exceptional downstream processing speed.

4.2.2. Task 3: Clustering

As shown in Table 4, the GPLSI submissions for Task 3 show that performance depends strongly on both the candidate-generation strategy and the target number of centroids k . Overall, our methods improve cluster cohesion with respect to the official baseline, measured by DBI, although their impact on retrieval effectiveness, measured by $nDCG@10$, varies across configurations. Among the three final submitted GPLSI approaches, *QMGC* is the most robust method, since it is the only strategy that consistently improves the baseline in both metrics across the three tested values of k .

For $k = 10$, the best GPLSI result is obtained by 10_SA_gplsi_QMGC005, with an $nDCG@10$ of 0.5628, slightly improving the baseline value of 0.5509. In terms of cluster compactness and separation, as measured by DBI, 10_QA_gplsi_QMGC006 reduces DBI from 7.9892 to 6.6033. Although GPLSI does not reach the best global $nDCG@10$ for this setting, 0.6363, it remains competitive in terms of DBI.

For $k = 25$, QMGC becomes more competitive. The submission 25_SA_gplsi_QMGC005 reaches an $nDCG@10$ of 0.5799, clearly improving the baseline value of 0.5284 and approaching the best global result of 0.5858. Similarly, 25_QA_gplsi_QMGC006 reduces DBI from 6.1201 to 5.1040, very close to the best global DBI of 5.0280.

The strongest GPLSI performance is obtained for $k = 50$. In this case, 50_QA_gplsi_QMGC006 achieves the best overall $nDCG@10$ among all submissions, with 0.5593, compared with the baseline value of 0.4656. Likewise, 50_QA_gplsi_QMGC005 obtains the best DBI, reducing it from 5.3679 to 4.2726. This shows that QMGC scales particularly well as the number of representative centroids increases.

Regarding the individual methods, *FPS* provides a purely geometric diversification strategy. It often improves DBI but does not consistently preserve retrieval quality. *QAL* adds query-awareness and slightly improves over FPS in $nDCG@10$, although its gains remain unstable. *QCSS*, described in the methodology as a spectral-seeding approach, was also evaluated with both SA and QA executions, but it was discarded during the internal screening phase used to select the five best results to be submitted, since the internal validation metrics ranked it as the weakest of the four explored approaches. A likely explanation is that QCSS mainly exploits the global spectral structure of the similarity graph, which may not be sufficiently aligned with the query-specific relevance signals required by this task. As a result, its selected centroids can be geometrically coherent but less effective at preserving retrieval-oriented

relevance, leading to a weaker balance between DBI and $nDCG@10$. In contrast, *QMGC* combines query-awareness with graph-based structural reduction, which explains its stronger balance between ranking preservation and cluster compactness.

Averaging across the three values of k and both execution variants, *FPS* obtains an average $nDCG@10$ of 0.4578 and a DBI of 5.9148, while *QAL* reaches 0.4776 and 5.9301. Both improve the average baseline DBI of 6.4924 but remain below the average baseline $nDCG@10$ of 0.5150. By contrast, *QMGC* achieves the best overall balance, with an average $nDCG@10$ of 0.5455 and the lowest average DBI among GPLSI methods, 5.4606.

Finally, *QA* executions show a clear runtime advantage over *SA*. While *SA* remains competitive in $nDCG@10$, especially for $k = 10$ and $k = 25$, *QA* obtains comparable or better DBI values with much lower annealing times. Overall, the results suggest that the best trade-off is achieved by *QMGC*, especially when combined with *QA* execution.

5. Conclusions

This year *Quantum CLEF Lab*, our team GPLSI presented multiple solutions for both instance selection (Task 2) and clustering tasks (Task 3). For Task 2, we introduced two different approaches, a *SVM augmented LocalSets*-based approach, achieving the second best result in terms of F_1 but having the drawback of lowest average reduction of the dataset size. While our *Stratified Partitioning RL* method achieved the highest reduction rate among submissions, it consequently resulted in the lowest F_1 score. The results obtained showed a notable trade-off between the number of selected instances and model performance.

For Task 3, we evaluated several candidate-generation strategies *FPS*, *QAL*, *QCSS* across various numbers of centroids (k). Our results showed that while most of our proposed methods successfully improve cluster cohesion (DBI) relative to the baseline, their impact on retrieval effectiveness ($nDCG@10$) varied significantly. The *QMGC* method emerged as our most robust and successful strategy, improving the metrics across all tested values of k and scaled exceptionally well, achieving the best overall $nDCG@10$ and DBI among all lab submissions for $k = 50$.

In general, our results demonstrate the crucial balance between data reduction and model performance. Across both tasks, we show that with the appropriate strategy, it is feasible to significantly reduce the amount of training data or clustering complexity while minimizing the impact on task performance.

Key contributions of this work include the following:

- **Multi-Paradigm methodologies for Instance Selection:** We introduce two complementary approaches that successfully navigate the trade-off between classification fidelity and data compression. These include a precision-driven, SVM-augmented *LocalSets* method that preserves boundary-critical instances, and a compression-driven Reinforcement Learning DQN approach based on stratified topological partitioning that achieves extreme dataset reduction.
- **Unified Quantum-Inspired Clustering Architecture:** We propose a modular clustering pipeline designed specifically for quantum and simulated annealers. We design and evaluate four distinct pivot selection strategies, highlighting the novel Query-aware Multilevel Graph Clustering *QMGC* method, which injects query-aligned structural signals to consistently balance ranking preservation and cluster compactness.

Future work will focus on addressing the observed limitations by exploring alternative configurations of our proposed methods. For Task 2, we plan to modify the second phase of the SVM-augmented *LocalSets* approach to achieve a more aggressive dataset reduction. Furthermore, for the *Stratified Partitioning RL* method, we aim to implement alternative clustering algorithms and reduce the total number of clusters (K) to increase the instance density per partition. We also intend to transition to a cosine distance-based QUBO formulation, as it more accurately captures semantic similarities within text embeddings. Finally, for Task 3, we plan to refine the *QMGC* method to maintain stable and balanced performance across varying target centroid sizes (k).

Acknowledgments

Funding: This research has been supported by the FPI grant (PREP2024-003110) from the Spanish Ministry of Science, Innovation and Universities, under the project “QUMLAUDE: Quantum Mechanics for Language Understanding and generation” (PID2024-160791OB-I00), funded by MCIN/AEI/10.13039/501100011033/ and by FEDER, UE; and project “SAFEWORDS: Language Anonymization with Ethical and Legal Safeguards through NLP” (AIA2025-163322-C63) funded by Ministry of Science, Innovation and Universities of the Spanish Government. This work is also partially funded by the Ministerio para la Transformación Digital y de la Función Pública and Plan de Recuperación, Transformación y Resiliencia - Funded by EU – NextGenerationEU within the framework of the project Desarrollo Modelos ALIA; and by the Ramón Areces Foundation through the project Criteria for Evaluation of Quality Corpora in Artificial Intelligence (CRITERIA), developed within the framework of the II National Call for Research Grants in the Humanities 2025, under the theme “Digital Humanities” (Reference: FRAHUMANIDADES25-01); Jairo Madrigal-Cutillas has a collaboration grant (num. application 25CO1/000382) funded by the Ministry of Education, Professional Training and Sports; Mónica Sánchez-Bellido has a research initiation grant (num. application 2025AIS380974) funded by the Vicerrectorado de Investigación of the Universidad de Alicante.

Declaration on Generative AI

During the preparation of this work, the author(s) used ChatGPT, Gemini, Claude in order to: Grammar, spelling check and rephrasing. After using these tools/services, the author(s) reviewed and edited the content as needed and take(s) full responsibility for the publication’s content.

References

- [1] A. Pasin, M. F. Dacrema, W. Cunha, M. A. Gonçalves, P. Cremonesi, N. Ferro, Quantumclef 2026: Overview of the third quantum computing challenge for information retrieval and recommender systems at CLEF, in: Working Notes of the Conference and Labs of the Evaluation Forum, CLEF 2026, Jena, Germany, 21-24 September 2026, CEUR Workshop Proceedings, CEUR-WS.org, 2026.
- [2] A. Pasin, M. F. Dacrema, W. Cunha, M. A. Gonçalves, P. Cremonesi, N. Ferro, Overview of quantumclef 2026: The third quantum computing challenge for information retrieval and recommender systems at CLEF, in: Experimental IR Meets Multilinguality, Multimodality, and Interaction - 17th International Conference of the CLEF Association, CLEF 2026, Jena, Germany, September 21-24, 2026, Proceedings, Lecture Notes in Computer Science, Springer, 2026.
- [3] J. P. Consuegra-Ayala, A. Morote-Martínez, F. Valero-Abellón, E. Lloret, P. Moreda, M. Palomar, Team gplsi at qclef 2025: Quantum-inspired instance selection and clustering, Working Notes of CLEF (2025).
- [4] C. Hutto, E. Gilbert, Vader: A parsimonious rule-based model for sentiment analysis of social media text, Proceedings of the International AAAI Conference on Web and Social Media 8 (2014) 216–225. URL: <https://ojs.aaai.org/index.php/ICWSM/article/view/14550>. doi:10.1609/icws.v8i1.14550.
- [5] A. Pasin, M. Ferrari Dacrema, W. Cunha, M. A. Gonçalves, P. Cremonesi, N. Ferro, et al., Quantumclef 2025: Overview of the second quantum computing challenge for information retrieval and recommender systems at clef, in: CEUR WORKSHOP PROCEEDINGS, volume 4038, CEUR-WS, 2025, pp. 4086–4104.
- [6] E. Leyva, A. González, R. Pérez, Three new instance selection methods based on local sets: A comparative study with several approaches from a bi-objective perspective, Pattern Recognition 48 (2015) 1523–1537. URL: <https://www.sciencedirect.com/science/article/pii/S0031320314003859>. doi:<https://doi.org/10.1016/j.patcog.2014.10.001>.

- [7] Y. Caises, A. González, E. Leyva, R. Pérez, Combining instance selection methods based on data characterization: An approach to increase their effectiveness, *Information Sciences* 181 (2011) 4780–4798. URL: <https://www.sciencedirect.com/science/article/pii/S0020025511002878>. doi:<https://doi.org/10.1016/j.ins.2011.06.013>, special Issue on Interpretable Fuzzy Systems.
- [8] A. Jha, H. Gupta, A. Nandi, RL-guided data selection for language model finetuning, 2025. URL: <https://arxiv.org/abs/2509.25850>. arXiv:2509.25850.
- [9] V. Mnih, K. Kavukcuoglu, D. Silver, A. Graves, I. Antonoglou, D. Wierstra, M. Riedmiller, Playing atari with deep reinforcement learning, 2013. URL: <https://arxiv.org/abs/1312.5602>. arXiv:1312.5602.
- [10] C. Bauckhage, R. Sifa, S. Wrobel, Adiabatic quantum computing for max-sum diversification, 2020. URL: <https://publica.fraunhofer.de/handle/publica/408331>. doi:10.1137/1.9781611976236.39.
- [11] C. Bauckhage, N. Piatkowski, R. Sifa, D. Hecker, S. Wrobel, A qubo formulation of the k-medoids problem, in: *Lernen, Wissen, Daten, Analysen (LWDA 2019)*, volume 2454 of *CEUR Workshop Proceedings*, 2019, pp. 54–63. URL: https://ceur-ws.org/Vol-2454/paper_39.pdf.
- [12] T. F. Gonzalez, Clustering to minimize the maximum intercluster distance, *Theoretical Computer Science* 38 (1985) 293–306. doi:10.1016/0304-3975(85)90224-5.
- [13] J. G. Carbonell, J. Goldstein, The use of mmr, diversity-based reranking for reordering documents and producing summaries, in: *Proceedings of the 21st Annual International ACM SIGIR Conference on Research and Development in Information Retrieval (SIGIR '98)*, 1998, pp. 335–336. URL: https://ir.webis.de/anthology/1998.sigirconf_conference-98.43/. doi:10.1145/290941.291025.
- [14] A. Loukas, Graph reduction with spectral and cut guarantees, *Journal of Machine Learning Research* 20 (2019) 1–42. URL: <https://jmlr.csail.mit.edu/papers/v20/18-680.html>.
- [15] D. Sculley, Web-scale k-means clustering, in: *Proceedings of the 19th International Conference on World Wide Web (WWW)*, 2010, pp. 1177–1178. doi:10.1145/1772690.1772862.
- [16] J. Shi, J. Malik, Normalized cuts and image segmentation, *IEEE Transactions on Pattern Analysis and Machine Intelligence* 22 (2000) 888–905. URL: <https://publications.ri.cmu.edu/normalized-cuts-and-image-segmentation>. doi:10.1109/34.868688.
- [17] A. Y. Ng, M. I. Jordan, Y. Weiss, On spectral clustering: Analysis and an algorithm, in: *Advances in Neural Information Processing Systems 14 (NeurIPS 2001)*, 2002, pp. 849–856. URL: https://proceedings.neurips.cc/paper_files/paper/2001/file/801272ee79cfde7fa5960571fee36b9b-Paper.pdf.
- [18] U. von Luxburg, A tutorial on spectral clustering, *Statistics and Computing* 17 (2007) 395–416. doi:10.1007/s11222-007-9033-z.