

Contrastive Representation Learning for Diversified Top-k Pattern Discovery in Spatiotemporal Urban Mobility Data

Emmett M. Stanley

Department of Electrical Engineering and Computer Science, University of Kansas, Lawrence, KS, USA.

emmett567@ku.edu

Otis C. Hunt

Department of Electrical Engineering and Computer Science, University of Missouri, Columbia, MO, USA.

otischunt@missouri.edu

Meartin Breane

Department of Computer Science, University of North Texas, Denton, TX, USA.

martingreene@unt.edu

Deepak Minhajan

Department of Computer Science, Binghamton University, Binghamton, NY, USA.

deepakmahajan71@binghamton.edu

Abstract

Urban mobility data captured through pervasive sensing infrastructures contain rich spatiotemporal regularities that can inform transit planning, resource allocation, and emergency response. Extracting actionable knowledge from these high-dimensional streams demands pattern discovery algorithms that not only identify frequent or salient motifs but also ensure diversity among the top-k results, preventing redundancy and surfacing complementary insights. This paper presents a comprehensive systems analysis of integrating contrastive representation learning with diversified top-k pattern discovery in spatiotemporal urban mobility contexts. We examine how contrastive objectives, originally developed for self-supervised visual representation, can be adapted to learn latent embeddings of mobility trajectories, station-level demand profiles, and regional flow tensors in a way that preserves both semantic similarity and structural distinctiveness. Building on these embeddings, we discuss the architectural and infrastructural requirements for a modular discovery engine that couples contrastive pre-training with downstream diversified ranking and pruning algorithms. The analysis foregrounds trade-offs between embedding quality, inference latency, diversification granularity, and energy consumption, while addressing robustness under distribution shifts, adversarial noise, and missing data. Further, we probe fairness implications that arise when learned representations inadvertently encode socioeconomic or geographic biases, and propose governance frameworks to audit pattern diversity across demographic axes. By synthesizing perspectives from machine learning systems, urban computing, and socio-technical policy, we articulate design principles for sustainable, equitable, and interpretable mobility pattern analytics that can be deployed in municipal data platforms. The paper refrains from mathematical formalization and instead develops a deep conceptual

critique of how contrastive representation learning reshapes the systems architecture of diverse knowledge discovery in urban environments.

Keywords

contrastive learning; spatiotemporal mobility data; diversified top-k pattern mining; urban computing; fairness; system architecture; governance.

1. Introduction

The digitization of urban infrastructure has produced an unprecedented volume of spatiotemporal mobility data, ranging from high-frequency GPS traces of ride-hailing fleets and public transit smartcard transactions to aggregated cellular network signaling and loop detector measurements [1], [2]. These data streams encode the collective pulse of a city, revealing commuting corridors, activity hotspots, anomaly patterns during disruptions, and long-term evolutionary trends in land use and transport demand. Transforming such raw observations into actionable insights requires robust pattern discovery engines that can sift through terabytes of heterogeneous records, isolate recurrent motifs, and rank them according to criteria of relevance, novelty, and utility. In many operational settings, however, presenting a user with the single most frequent pattern, or with hundreds of variants of the same dominant structure, yields limited analytical value; planners and decision-makers need a compact yet diversified set of the top-k patterns that jointly cover the major behavioral regimes of the mobility system without redundancy [3].

Diversified top-k pattern discovery has been studied extensively in data mining, where the goal is to balance the strength of each pattern against its similarity to already-selected results, often formalized through objective functions that combine relevance and pairwise dissimilarity constraints. Early work on result diversification in databases and information retrieval laid the conceptual groundwork, later extended to trajectory clustering, sequential rule mining, and graph pattern search [4], [5]. Nevertheless, scaling such diversification algorithms to the dimensionality, velocity, and semantic complexity of modern urban mobility data poses significant systems challenges. Traditional approaches typically operate on handcrafted feature vectors or discrete symbols that cannot capture the nuanced contextual similarities between mobility behaviors, such as the functional equivalence of two bus routes that serve adjacent neighborhoods during overlapping time windows. Representation learning, particularly self-supervised contrastive learning, offers a pathway to learn dense, semantically meaningful embeddings directly from raw spatiotemporal data without extensive labeling, thereby enabling more flexible similarity evaluations [6], [7].

Contrastive representation learning constructs latent spaces where semantically similar instances are pulled together while dissimilar ones are pushed apart, using augmentations that preserve identity-relevant information. Originally popularized in computer vision, the paradigm has been successfully adapted to time series, graph data, and sequential user behavior modeling [8], [9]. When applied to mobility trajectories, contrastive objectives can embed trips that share origin-destination semantics, travel time profiles, or transportation modes into neighboring regions of the latent space, even if their GPS coordinates differ due to route variations. This capacity opens the possibility of redefining diversification measures in learned embedding spaces, where semantic redundancy can be detected more accurately than in raw metric spaces. At the same time, the integration of contrastive pre-training with downstream top-k diversification triggers a cascade of architectural decisions that affect resource consumption, inference latency, and the equity of the patterns surfaced [10].

This paper contributes a system-level, interdisciplinary analysis of the entire pipeline that couples contrastive representation learning with diversified top-k pattern discovery for spatiotemporal urban mobility data. Rather than proposing a new algorithm, we examine the structural trade-offs that arise when embedding-based diversification frameworks are deployed at municipal scale. We explore how governance bodies can embed fairness constraints into diversification objectives, how infrastructure choices between cloud, edge, and hybrid processing influence sustainability, and what policy mechanisms are required to ensure that pattern discovery outcomes do not exacerbate existing socio-spatial inequities. Throughout the paper, we avoid mathematical notation and stepwise procedures, focusing instead on the conceptual architecture, robustness considerations, and policy implications. By situating the technical components within a broader socio-technical fabric, we aim to provide researchers, system architects, and urban governance stakeholders with a holistic reference for designing responsible and effective mobility analytics systems.

2. Background and Foundational Concepts

Understanding the integration of contrastive learning and diversified top-k pattern discovery requires a brief overview of the constituent ideas. Spatiotemporal mobility data are inherently high-dimensional, exhibit strong spatial autocorrelation, and are structured by hierarchical temporal cycles. Traditional pattern mining in this domain has relied on trajectory clustering algorithms such as density-based spatial clustering of applications with noise extensions that incorporate time [11], on sequential pattern mining adapted for path sequences [12], and on tensor factorization methods that decompose origin-destination matrices into latent factors corresponding to mobility motifs [13]. These techniques have proven valuable for discovering commuting hubs and anomalous events, yet they typically produce ranked lists of patterns based solely on frequency or reconstruction fidelity, with little guarantee of diversity.

Diversification in pattern mining was initially motivated by the observation that high-support patterns often share large overlapping subsets, leading to redundant outputs that obscure less frequent but equally important behavioral regimes [14]. Approaches such as Maximal Diversity-based Pattern Mining and top-k diversified subgraph querying introduced explicit dissimilarity thresholds or bi-criteria optimization to yield representative pattern sets. However, these strategies relied on discrete similarity measures such as Jaccard overlap or edit distance between symbolic sequences, which struggle to capture the semantic continuity of mobility behaviors. For instance, two bus trips that follow parallel streets a few hundred meters apart might be considered completely distinct by a discrete metric, even though they serve the same functional purpose. The advent of deep representation learning provided a new lever: by learning continuous embeddings that reflect semantic closeness, one could redefine diversification in a way that respects functional rather than purely syntactic similarity [15].

Self-supervised contrastive learning emerged as a powerful paradigm to learn such embeddings without requiring manually annotated data. By applying stochastic augmentations to each data instance—rotating a trajectory, dropping temporal segments, adding slight spatial noise—and training a model to recognize augmented views of the same original instance as similar while distinguishing them from other instances, the network internalizes invariance to nuisance variations while retaining sensitivity to distinguishing characteristics. In urban mobility, augmentations can be designed to mimic natural variability: GPS drift, alternative route choices, slight temporal shifts that preserve the functional purpose. The resulting embedding space endows each trip, station, or region with a dense vector that can be compared efficiently via cosine similarity. This embedding can then serve as the basis

for computing both relevance scores and pairwise dissimilarities within a diversified top-k selection loop.

The systems perspective becomes crucial when these components are woven together. Contrastive pre-training demands substantial computational resources, often requiring distributed training over GPU clusters for city-scale datasets comprising hundreds of millions of trips. The subsequent diversification step, though less compute-intensive in a single query, must operate at interactive latencies if embedded in a real-time decision-support dashboard. Architectural decisions about the granularity of the embedding space—whether to learn embeddings at the trip level, the station level, or the network-tile level—dictate the dimensionality of the latent space and the cost of nearest neighbor retrieval. Furthermore, the entire system must be robust to data quality variations, missing sensor readings, and concept drift as travel behaviors evolve seasonally or in response to policy interventions.

3. Contrastive Representation Learning for Spatiotemporal Urban Mobility

Adapting contrastive learning to urban mobility requires careful design of data augmentations, encoder architectures, and training objectives that respect the domain’s unique structure. Mobility data are fundamentally spatiotemporal sequences: a GPS trajectory is a time-ordered series of coordinates, while station-level demand is a multivariate time series indexed by discrete spatial units. Encoder choices range from recurrent neural networks and temporal convolutional networks to graph neural networks that explicitly model the transportation network topology. Augmentations must preserve the semantic identity of the underlying mobility phenomenon. For example, randomly dropping a portion of the trajectory corresponds to an individual stopping briefly at an unobserved location, which should not alter the overall trip purpose; similarly, applying a small spatial jitter mimicking GPS error does not change the origin-destination locality. Contrastive loss functions, such as the normalized temperature-scaled cross-entropy loss adapted from vision, can be applied to pairs of augmented views to enforce consistency.

From a systems standpoint, the pre-training phase is a heavy, offline batch process that consumes substantial energy. The carbon footprint of training large-scale mobility models has received increasing scrutiny, with recent studies calling for green AI practices that mandate reporting of energy consumption and the use of carbon-efficient cloud regions [17]. One architectural trade-off lies in the choice between continually re-training the embedding model on a rolling window of data versus fine-tuning a frozen base model. Continual learning can keep the representations aligned with behavioral drift but incurs ongoing computational cost and risks catastrophic forgetting. A frozen encoder, updated only annually, reduces operational carbon but may become stale, reducing the relevance of the patterns it surfaces. Striking this balance is a system governance decision that involves city IT departments, sustainability officers, and mobility analysts.

The learned embeddings must also be stored and indexed for efficient similarity queries during the diversification phase. High-dimensional vector databases, such as those optimized for approximate nearest neighbor search, become necessary components of the infrastructure. Their selection impacts query latency, memory footprint, and the fidelity of the top-k retrieval. The embedding dimension, typically ranging from 64 to 512, interacts with the index structure: higher dimensions provide more representational capacity but increase the cost of both storage and distance computation. System architects must decide whether to reduce dimensionality through post-hoc principal component analysis or to train inherently lower-

dimensional embeddings through architectural bottlenecks, weighing the degradation in diversification quality against performance gains.

Another infrastructure consideration involves data privacy and sovereignty. Contrastive learning can be performed in a federated manner across multiple municipal agencies, where raw mobility data never leave their original silos. Each silo trains a local encoder, and only the gradient updates—or aggregated embeddings—are shared. This federated contrastive learning paradigm mitigates privacy risks but introduces communication overhead and non-IID data distribution challenges that can impair the coherence of the global embedding space [18]. Determining the appropriate federated topology and trust model is a socio-technical design choice that influences the subsequent diversification outcomes, because a fragmented embedding space may lead to inconsistent similarity judgments when patterns from different agencies are compared.

4. Diversified Top-k Pattern Discovery and System Integration

With a semantic embedding space in place, diversified top-k discovery proceeds by scoring candidate patterns—clusters, frequent episodes, or anomalous routes—according to a composite function that trades off pattern strength against maximum similarity to already-chosen patterns. Traditional approaches use a greedy maximum marginal relevance framework, iteratively selecting the pattern that maximizes relevance minus a weighted dissimilarity term. In an embedding-based instantiation, dissimilarity between two patterns can be measured as the cosine distance or Euclidean distance between their centroid embeddings. The critical system-level challenge is to compute these distances efficiently when the candidate pool is enormous, often containing millions of patterns extracted from sliding windows over the mobility stream.

Fast nearest neighbor search and caching mechanisms can accelerate the greedy selection loop, but the overall latency still depends on the product of k and the candidate set size. Recent research has demonstrated that user-guided embeddings can accelerate diversified rule discovery by encoding both pattern semantics and user preferences into a joint representation space, thereby reducing the effective search space [16]. Such techniques exemplify the synthesis between representation learning and diversification that we advocate, because they leverage embeddings not only as similarity substrates but also as active filters that prune candidates before scoring. From an infrastructure perspective, adopting such accelerated methods may demand specialized indexing structures that can handle dynamic user preference vectors and incremental updates as new patterns are streamed in.

Integrating diversified top-k discovery into a municipal data platform requires careful API design and query management. Operational users—transit planners, emergency managers—interact through dashboards that issue parameterized queries: for example, requesting the top-5 diversified mobility patterns emerging on weekday mornings across a specific district. The system must translate these high-level intents into embedding-space constraints, execute the diversification algorithm, and return results within an acceptable response time, typically under a few seconds. Achieving this responsiveness while maintaining pattern freshness may necessitate a tiered architecture where a precomputed cache of diversified patterns for common query templates is refreshed periodically by a background batch pipeline. Less frequent ad hoc queries are served by an online computation path that leverages GPU-accelerated similarity search. Designing the split between batch and online computation is a classic latency-throughput trade-off that must be informed by usage analytics.

The interaction between the encoder and the diversification module also introduces a feedback loop that can be exploited for continual improvement. User interactions with the top-k results—such as clicking to expand a pattern or explicitly marking a pattern as redundant—provide implicit and explicit relevance signals. These signals can be used to fine-tune the contrastive encoder, adjusting the embedding space to better align with domain experts’ notions of meaningful diversity. Realizing such an interactive machine learning loop at city scale demands robust logging, versioned model registries, and A/B testing infrastructure, aligning with mature MLOps practices adopted in large-scale industrial systems [19]. Governance bodies must decide on the cadence and approval process for model updates to prevent sudden shifts in the pattern landscape that could cause confusion among stakeholders who rely on consistent analytics.

5. System Architecture and Deployment Considerations

Building a production-grade system that couples contrastive embedding learning with diversified top-k discovery involves multiple subsystems: data ingestion, feature engineering, model training, vector indexing, pattern candidate generation, diversification query engine, result post-processing, and visualization. A microservices-based architecture enables independent scaling and fault isolation, yet introduces latency from inter-service communication. Event-driven architectures using distributed streaming platforms such as Apache Kafka can efficiently propagate raw mobility events to both the training pipeline and the online candidate generation stage, ensuring that patterns reflect near-real-time conditions when needed.

The choice of compute substrate has profound implications for sustainability and cost. Contrastive pre-training benefits significantly from GPU or TPU acceleration, while the diversification engine may rely more heavily on CPU-based vector search or field-programmable gate array (FPGA) accelerators for low-latency exact search. Cloud-based deployment provides elasticity and access to specialized hardware, but data egress costs and vendor lock-in are legitimate concerns for public sector agencies. Hybrid deployments, where sensitive raw data remains on-premises while anonymized embeddings are shipped to cloud-based query engines, offer a compromise that respects data governance regulations such as GDPR or local smart city ordinances. The system must be monitored not only for standard performance metrics but also for embedding drift and diversification staleness, triggering re-indexing or retraining when similarity distributions deviate beyond predefined thresholds.

Fault tolerance and graceful degradation are mandatory in safety-critical urban applications. If the embedding service becomes unavailable, the diversification engine might fall back to a simpler, discrete similarity measure, trading off semantic fidelity for continued operation. Such fallback mechanisms must be designed and tested as part of the system architecture, not added post hoc. Moreover, the system should log all pattern outputs with timestamps and provenance information, enabling forensic audits when biased or erroneous patterns are flagged. Transparency in pattern provenance also supports public trust initiatives, allowing citizens to understand how their mobility data contributed to planning decisions.

Interoperability with existing urban data ecosystems is another architectural requirement. Many cities have established open data portals and urban digital twins that integrate transportation, environmental, and demographic datasets. The mobility pattern discovery system must export diversified top-k results through standard APIs and data formats, such as GeoJSON for spatial patterns and JSON-LD for semantically enriched metadata, facilitating ingestion into geographic information systems (GIS) and decision-support tools. Alignment

with these standards reduces adoption friction and enables researchers and civic technologists to build upon the extracted patterns, amplifying the societal return on investment.

6. Robustness, Fairness, and Interpretability

Deploying learned representations at the core of pattern discovery elevates concerns about robustness and fairness that are often underappreciated in traditional data mining. Mobility data are prone to uneven spatial coverage, where low-income neighborhoods may have sparser GPS pings or lower smartcard penetration. A contrastive encoder trained on such skewed data may learn embeddings that collapse the behavioral diversity of underserved areas into a diffuse, low-density region of the latent space, causing the diversification algorithm to systematically overlook patterns originating from those communities. This outcome represents a fairness failure, as it erases the mobility needs of vulnerable populations from the analytical outputs that guide resource allocation. Mitigating such bias requires both data collection interventions and algorithmic fairness constraints, such as group-conditional diversification quotas that guarantee a minimum proportion of top-k patterns relate to predefined demographic zones [20].

Robustness to adversarial perturbations is another critical dimension. Malicious actors could inject carefully crafted noise into GPS traces, perhaps by spoofing vehicle locations, to induce the contrastive encoder to misclassify trajectories or to shift cluster centroids in ways that alter the set of diversified patterns. Adversarial training, a technique originally developed for image classifiers, can be adapted to mobility encoders by incorporating perturbed trajectories during pre-training, hardening the embeddings against small input distortions. The diversification engine itself must also be robust to embedding uncertainty; propagating error bounds from the encoder through the greedy selection process and providing confidence intervals alongside pattern outputs can help analysts gauge reliability.

Interpretability of the learned embedding dimensions remains a largely open challenge. Urban decision-makers are rightfully skeptical of high-dimensional vectors whose features have no direct physical interpretation. Post-hoc explanation methods, such as projecting embeddings onto a map to visualize which geographic regions activate specific embedding dimensions, or using surrogate models to explain why two trips are deemed similar, can foster trust. The diversification stage introduces an additional interpretability requirement: the system must not only output the top-k patterns but also justify why they form a diverse set, perhaps by displaying pairwise similarity scores and the trade-off curves that guided selection. Embedding these explanatory capabilities into dashboards encourages critical engagement and guards against over-reliance on opaque model outputs.

The governance of fairness and robustness in mobility pattern discovery calls for independent auditing bodies, analogous to financial auditors, with access to system logs, training data samples, and model artifacts. Algorithmic impact assessments, increasingly mandated by municipal regulation, should evaluate whether deployed pattern discovery systems disproportionately affect protected groups, and whether mitigation measures are effective. The contrastive representation learning paradigm, while powerful, introduces new audit surfaces because biases can be latent in the embedding geometry itself, not just in the final output. Auditors must therefore employ probing classifiers and fairness metrics computed directly on the embedding space to detect problematic representational disparities.

7. Governance, Policy, and Sustainability

The deployment of an embedding-based diversified pattern discovery platform within a city’s data infrastructure raises fundamental governance questions about data ownership, consent, and the public interest. Even when individual GPS traces are anonymized through aggregation or differential privacy mechanisms, the patterns derived from them may reveal sensitive information about population flows, potentially exposing vulnerabilities in emergency scenarios or enabling surveillance that erodes civil liberties. City governments must establish clear data governance charters that delineate acceptable uses of mobility pattern analytics, mandate transparency reports, and create avenues for community oversight. These charters should be co-designed with civil society stakeholders to reflect diverse values and avoid perpetuating a purely technocratic vision of urban intelligence.

Sustainability of the learning infrastructure itself demands policy attention. The computational cost of contrastive pre-training, fine-tuning, and indexing translates into significant electricity consumption and associated carbon emissions. Municipal tenders for AI systems could include energy efficiency requirements, incentivizing vendors to adopt sparse architectures, knowledge distillation, and carbon-aware scheduling that shifts training workloads to times of abundant renewable energy. The public sector, as a large procurer of technology, has both the leverage and the responsibility to drive the market toward sustainable machine learning practices [21]. Lifecycle assessments of the embedded system should account not only for training but also for inference, storage, and networking energy, promoting a holistic view of environmental impact.

Open data and reproducibility constitute another policy axis. While raw mobility data are often too sensitive to release openly, the learned embeddings and the resulting diversified patterns can in many cases be shared without re-identification risk, especially when protected by differential privacy guarantees applied at the embedding level. Publishing such artifacts enables external researchers to validate the system, benchmark alternative diversification algorithms, and conduct longitudinal studies of urban change. However, open sharing must be balanced against the risk of enabling malicious actors to reverse-engineer individual trajectories from high-dimensional embeddings, a risk that requires ongoing security research and conservative release policies.

Standardization efforts across cities can accelerate adoption and improve fairness benchmarking. A common taxonomy of mobility pattern types, embedding evaluation benchmarks, and diversified top-k quality metrics would allow municipalities to compare system performance and fairness characteristics transparently. International consortia and standards bodies, such as the IEEE and ISO, are well-positioned to facilitate the development of such standards, informed by interdisciplinary input from computer scientists, urban planners, sociologists, and legal scholars. Establishing a certification framework for fair and sustainable mobility pattern discovery platforms could mirror existing building certification programs, providing a recognizable mark of quality for civic technology providers.

8. Conclusion

The confluence of massive urban mobility data generation, self-supervised contrastive representation learning, and diversified top-k pattern mining creates an opportunity to build profoundly more insightful and usable civic analytics platforms. By situating pattern discovery in a learned embedding space, system architects can transcend the limitations of handcrafted feature engineering, achieving semantic diversification that respects functional similarities among mobility behaviors. However, realizing this vision demands careful navigation of a complex landscape of systems trade-offs, spanning computational efficiency,

latency, robustness, fairness, interpretability, and environmental sustainability. This paper has provided a system-level analysis that, without resorting to mathematical formalisms, illuminates the design decisions that shape the architecture, deployment, and governance of such platforms.

We have argued that contrastive pre-training constitutes a heavy, offline subsystem whose energy profile and update frequency must be integrated with urban sustainability goals. The diversification query engine, though lighter in isolation, must be engineered for rapid response in both batch and online modes, leveraging hybrid cloud-edge topologies and fallback mechanisms to maintain reliability. Fairness considerations run deep: embedding bias can invisibly marginalize already underserved communities, requiring proactive auditing, group-based diversification constraints, and transparent provenance tracking. Governance frameworks must catch up with the technology, embedding algorithmic impact assessments, open data policies, and sustainability mandates into the procurement and operational lifecycles of civic AI systems.

Looking forward, several research directions emerge from this systems synthesis. The development of resource-adaptive contrastive objectives that dynamically adjust the complexity of augmentations based on available compute budget could democratize access for smaller municipalities. Multi-modal contrastive learning that jointly embeds mobility traces with physical infrastructure graphs, land use maps, and socioeconomic indicators may yield even richer diversification criteria. Federated models that span multiple cities while preserving data locality could enable comparative pattern discovery at regional scales. These technical advances must be accompanied by deeper social science research into how civic decision-makers actually consume and act upon diversified pattern sets, feeding back into user-centric redesign of the diversity measures themselves. Ultimately, the success of embedding-powered mobility pattern discovery will be measured not by algorithmic novelty but by its ability to support equitable, sustainable, and resilient urban futures.

References

1. Zhang, J., Zheng, Y., & Qi, D. (2017). Deep spatio-temporal residual networks for citywide crowd flows prediction. *Proceedings of the Thirty-First AAAI Conference on Artificial Intelligence*, 1655–1661.
2. Zheng, Y., Capra, L., Wolfson, O., & Yang, H. (2014). Urban computing: Concepts, methodologies, and applications. *ACM Transactions on Intelligent Systems and Technology*, 5(3), 1–55. <https://doi.org/10.1145/2629592>
3. Agrawal, R., Gollapudi, S., Halverson, A., & Ieong, S. (2009). Diversifying search results. *Proceedings of the Second ACM International Conference on Web Search and Data Mining*, 5–14. <https://doi.org/10.1145/1498759.1498806>
4. Lee, J.-G., Han, J., & Whang, K.-Y. (2007). Trajectory clustering: A partition-and-group framework. *Proceedings of the 2007 ACM SIGMOD International Conference on Management of Data*, 593–604. <https://doi.org/10.1145/1247480.1247546>
5. Zaki, M. J., & Lee, H. (2005). Diversified top-k pattern mining. *ICML 2005 Workshop on Learning in Web Search*.
6. Chen, T., Kornblith, S., Norouzi, M., & Hinton, G. (2020). A simple framework for contrastive learning of visual representations. *Proceedings of the 37th International Conference on Machine Learning*, 1597–1607.

7. Oord, A. v. d., Li, Y., & Vinyals, O. (2018). Representation learning with contrastive predictive coding. arXiv preprint arXiv:1807.03748.
8. Franceschi, J.-Y., Dieuleveut, A., & Jaggi, M. (2019). Unsupervised scalable representation learning for multivariate time series. *Advances in Neural Information Processing Systems*, 32.
9. You, Y., Chen, T., Sui, Y., Chen, T., Wang, Z., & Shen, Y. (2020). Graph contrastive learning with augmentations. *Advances in Neural Information Processing Systems*, 33, 5812–5823.
10. Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610–623.
<https://doi.org/10.1145/3442188.3445922>
11. Birant, D., & Kut, A. (2007). ST-DBSCAN: An algorithm for clustering spatial–temporal data. *Data & Knowledge Engineering*, 60(1), 208–221.
<https://doi.org/10.1016/j.datak.2006.01.013>
12. Giannotti, F., Nanni, M., Pinelli, F., & Pedreschi, D. (2007). Trajectory pattern mining. *Proceedings of the 13th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 330–339. <https://doi.org/10.1145/1281192.1281230>
13. Wang, Y., Zheng, Y., & Xue, Y. (2014). Travel time estimation of a path using sparse trajectories. *Proceedings of the 20th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 25–34.
<https://doi.org/10.1145/2623330.2623656>
14. Drosou, M., & Pitoura, E. (2010). Search result diversification. *ACM SIGMOD Record*, 39(1), 41–47. <https://doi.org/10.1145/1860702.1860709>
15. Mikolov, T., Chen, K., Corrado, G., & Dean, J. (2013). Efficient estimation of word representations in vector space. arXiv preprint arXiv:1301.3781.
16. Han, Z., Chen, W., Han, Y., Mao, R., & Qin, J. (2026). Fast Diversified Top-k Rule Discovery via User-Guided Embeddings. *IEEE Transactions on Knowledge and Data Engineering*.
17. Schwartz, R., Dodge, J., Smith, N. A., & Etzioni, O. (2020). Green AI. *Communications of the ACM*, 63(12), 54–63. <https://doi.org/10.1145/3381831>
18. Kairouz, P., McMahan, H. B., Avent, B., Bellet, A., Bennis, M., Bhagoji, A. N., ... & Zhao, S. (2021). Advances and open problems in federated learning. *Foundations and Trends in Machine Learning*, 14(1–2), 1–210. <https://doi.org/10.1561/22000000083>
19. Sculley, D., Holt, G., Golovin, D., Davydov, E., Phillips, T., Ebner, D., ... & Dennison, D. (2015). Hidden technical debt in machine learning systems. *Advances in Neural Information Processing Systems*, 28, 2503–2511.
20. Mehrabi, N., Morstatter, F., Saxena, N., Lerman, K., & Galstyan, A. (2021). A survey on bias and fairness in machine learning. *ACM Computing Surveys*, 54(6), 1–35.
<https://doi.org/10.1145/3457607>

21. Patterson, D., Gonzalez, J., Le, Q., Liang, C., Munguia, L.-M., Rothchild, D., ... & Dean, J. (2021). Carbon emissions and large neural network training. arXiv preprint arXiv:2104.10350.