



An Efficient and Fine-grained Big Data Access Control Scheme with Privacy-preserving Policy

PRAVEEN KUMAR

Bharati Vidyapeeth University, College Of Engineering, Pune

Abstract — How to management the access of the massive quantity of massive knowledge becomes a awfully difficult issue, particularly once huge knowledge are keep within the cloud. Cipher text-Policy Attribute primarily based coding (CP-ABE) may be a promising coding technique permits end-users to inscribe their knowledge below the access policies outlined over some attributes of knowledge customers and solely allows data customers whose attributes satisfy the access policies to rewrite the info[1]. In CP-ABE, the access policy is connected to the cipher text in plaintext type, which can additionally leak some personal info regarding end-users. Existing strategies solely partly hide the attribute values within the access policies, whereas the attribute names are still unprotected. During this paper, we have a tendency to propose Associate in Nursing economical and fine-grained huge knowledge access management theme with privacy-preserving policy[3]. Specifically, we have a tendency to hide the complete attribute (rather than solely its values) within the access policies. to help knowledge coding, we have a tendency to additionally style a unique Attribute Bloom Filter to gauge whether or not Associate in Nursing attribute is within the access policy and find the precise position within the access policy if it's within the access policy. Security analysis and performance analysis show that our theme will preserve the privacy from any LSSS access policy while not using abundant overhead.

Keywords- Big Data; Access Control; Privacy-preserving Policy; Attribute Bloom Filter; LSSS Access Structure.

I. INTRODUCTION

Big data is a term that refers to data sets or combinations of data sets whose size (volume), complexity (variability), and rate of growth (velocity) make them difficult to be captured, managed, processed or analyzed by conventional technologies and tools, such as relational databases and desktop statistics or visualization packages, within the time necessary to make them useful. While the size used to determine whether a particular data set is considered big data is not firmly defined and continues to change over time, most analysts and practitioners currently refer to data sets from 30-50 terabytes(10¹² or 1000 gigabytes per terabyte) to multiple petabytes (10¹⁵ or 1000 terabytes per petabyte) as big data.[4]

The analysis of Big Data involves multiple distinct phases as shown in the figure below, each of which introduces challenges. Many people unfortunately focus just on the analysis/modeling phase: while that phase is crucial, it is of little use without the other phases of the data analysis pipeline. Even in the analysis phase, which has received much attention, there are poorly understood complexities in the context of multi-tenanted clusters where several users' programs run concurrently. Many significant challenges extend beyond the analysis phase. For example, Big Data has to be managed in context, which may be noisy, heterogeneous and not include an upfront model. Doing so raises the need to track provenance and to handle uncertainty and error: topics that are crucial to success, and yet rarely mentioned in the same breath as Big Data. Similarly, the questions to the data analysis pipeline will typically not all be laid out in advance. It may need to figure out good questions based on the data. Doing this will require smarter systems and also better support for user interaction with the analysis pipeline. In fact, there is a major bottleneck in the number of people empowered to ask questions of the data and analyze it. It can drastically increase this number Big knowledge may be a term that refers to knowledge sets or mixtures of knowledge sets whose size (volume), quality (variability), and rate of growth (velocity) build them troublesome to be captured, managed, processed or analyzed by standard technologies and tools, like relative databases and desktop statistics or image packages, among the time necessary to form them helpful. whereas the dimensions accustomed verify whether or not a selected knowledge set is taken into account huge knowledge isn't firmly outlined and continues to vary over time, most analysts and practitioners presently talk to knowledge sets from 30-50 terabytes(10¹² or a thousand gigabytes per terabyte) to multiple petabytes (10¹⁵ or a thousand terabytes per petabyte) as huge knowledge.

The analysis of massive knowledge involves multiple distinct phases as shown within the figure below, every of that introduces challenges. many folks sadly focus simply on the analysis/modeling part: whereas that phase is crucial, it's of very little use while not the opposite phases of the info analysis pipeline. Even within the analysis part, that has received a lot of attention, there square measure poorly understood complexities within the context of multi-tenanted clusters wherever many users' programs run at the same time. several important challenges extend on the far side the analysis part. for instance, huge knowledge should be managed in context, which can be reedy, heterogeneous Associate in Nursingd not embody an direct model. Doing therefore raises the requirement to trace place of origin and to handle uncertainty and error: topics that square measure crucial to success, and nevertheless seldom mentioned within the same breath as huge knowledge. Similarly, the inquiries to the info analysis pipeline can generally not all be arranged move into advance. it should have to be compelled to understand smart queries supported the info. Doing this may need smarter systems and additionally higher support for user interaction with the analysis pipeline. In fact, there's a significant bottleneck within the range of individuals authorised to raise queries of the info and analyze it. It will drastically increase this range by supporting several levels of engagement with the info, not all requiring deep information experience. Solutions to issues like this may not come back from progressive enhancements to business as was common like trade may build on its own. fortuitously, existing process techniques may be applied, either as is or with some extensions, to a minimum of some aspects of the massive knowledge drawback. for instance, relative knowledge bases have faith in the notion of logical data independence: users will trust what they require to work out, whereas the system (with adept engineers coming up with those systems) determines the way to work out it with efficiency. Similarly, the SQL normal and also the relative knowledge model offer a consistent, powerful language to specific several question desires and, in essence, permits customers to decide on between vendors, increasing competition. The challenge previous Maine is to mix these healthy options of previous systems. Map scale back has emerged as a preferred thanks to harness the ability of huge clusters of computers. Map scale back permits programmers to assume in a very data-centric fashion: they concentrate on applying transformations to sets of knowledge records, and permit the main points of distributed execution, network communication and fault tolerance to be handled by the Map scale back framework. Map scale back is often applied to massive batch-oriented computations that square measure involved primarily with time to job completion. The Google Map scale back framework and ASCII text file Hadoop system reinforce this usage model through a batch-processing implementation strategy: the whole output of every map and scale back task is materialized to a neighborhood file before it may be consumed by subsequent stage. Materialization permits for an easy and chic checkpoint/restart fault tolerance mechanism that's vital in massive deployments, that have a high likelihood of slowdowns or failures at employee nodes .

II.LITERATURE SURVEY

1) Title: Privacy-Preserving Data Publishing: A Survey of Recent Developments

Written By: BENJAMIN C. M. FUNG, KE WANG, RUI CHEN, PHILIP S. YU.

The collection of digital info by governments, companies, and people has created tremendous opportunities for knowledge- and information-based higher cognitive process. Driven by mutual advantages, or by laws that need sure knowledge to be printed, there's a requirement for the exchange and publication of knowledge among varied parties. knowledge in its original type, however, usually contains sensitive info regarding people, and commercial enterprise such knowledge can violate individual privacy. this apply in knowledge commercial enterprise depends principally on policies and tips on what sorts of knowledge will be printed and on agreements on the utilization of printed knowledge. This approach alone might cause excessive knowledge distortion or short protection. Privacy-preserving knowledge commercial enterprise (PPDP) provides ways and tools for commercial enterprise helpful info whereas protective knowledge privacy. Recently, PPDP has received respectable attention in analysis communities, and plenty of approaches are projected for various knowledge commercial enterprise situations. during this survey, we'll consistently summarize and value totally different approaches to PPDP, study the challenges in sensible knowledge commercial enterprise, clarify the variations and needs that distinguish PPDP from different connected issues, and propose future analysis directions.

2) Title: APPLLET: a privacy-preserving framework for location-aware recommender system

Written By: Xindi Ma,HuiLI, Jianfeng MA, Qi JIANG, Sheng GAO, Ning XI &DiLU

Location-aware recommender systems that use location-based ratings to provide recommendations have recently intimate a speedy development and draw vital attention from the analysis community. However, current work principally centered on high-quality recommendations whereas underestimating privacy problems, which may cause issues of privacy. Such

issues are a lot of distinguished once service suppliers, WHO have restricted machine and storage resources, leverage on cloud platforms to suit in with the tremendous number of service necessities and users. During this paper, we have a tendency to propose a unique framework, specifically applications programmer, for shielding user privacy info, together with locations and recommendation results, inside a cloud surroundings. Through this framework, all historical ratings are hold on and calculated in cipher text, permitting US to firmly cipher the similarities of venues through Paillier secret writing, and predict the advice results supported Paillier, independent, and comparable secret writing. we have a tendency to conjointly on paper prove that user info is non-public and can not be leaked throughout a recommendation. Finally, empirical results over a real-world dataset demonstrate that our framework will with efficiency suggest POIs with a high degree of accuracy in a very privacy-preserving manner.

3) Title: Efficient Discovery of De-identification Policies Through a Risk-Utility Frontier

Written By: Weiyi Xia, Raymond Heatherly, Xiaofeng Ding

Modern data technologies modify organizations to capture giant quantities of person-specific knowledge whereas providing routine services. several organizations hope, or area unit de jure needed, to share such knowledge for secondary functions (e.g. validation of analysis findings) during a de-identified manner. In previous work, it had been shown de-identification policy alternatives might be sculpturesque on a lattice, that might be looked for policies that met a prespecified risk threshold (e.g., chance of re-identification). However, the search was restricted in many ways that. First, its definition of utility was syntactical supported the extent of the lattice - and not linguistics - based mostly on the particular changes evoked within the ensuing knowledge. Second, the edge might not be famous beforehand. The goal of this work is to create the optimum set of policies that trade-off between privacy risk (R) and utility (U), that we have a tendency to ask as a R-U frontier. To model this drawback, we have a tendency to in-troduce a linguistics definition of utility, supported scientific theory, that's compatible with the lattice illustration of policies. To unravel the matter, we have a tendency to at first build a group of policies that outline a frontier. We have a tendency to then use a probability-guided heuristic to go looking the lattice for policies possible to update the frontier. To demonstrate the effectiveness of our approach, we have a tendency to perform associate degree empirical analysis with the Adult dataset of the UCI Machine Learning Repository. We show that our approach will construct a frontier nearer to optimum than competitive approaches by looking a smaller range of policies. additionally, we have a tendency to show that a often followed de-identification policy (i.e., the porcupine provision customary of the HIPAA Privacy Rule) is suboptimal as compared to the frontier discovered by our approach.

4) Title: l -Diversity: Privacy Beyond k-Anonymity

Written By: Ashwin Machanavajjhala, Johannes Gehrke, Daniel Kifer

Publishing knowledge regarding people while not revealing sensitive data regarding them is a vital downside. In recent years, a replacement definition of privacy referred to as k-anonymity has gained quality. during a k-anonymized dataset, every record is indistinguishable from a minimum of k -1 different records with relation to bound "identifying" attributes. during this paper we have a tendency to show with 2 easy attacks that a k-anonymized dataset has some refined, however severe privacy problems. First, we have a tendency to show that AN offender will discover the values of sensitive attributes once there's very little diversity in those sensitive attributes. Second, attackers typically have background, and that we show that k-anonymity doesn't guarantee privacy against attackers mistreatment background. We have a tendency to provides a careful analysis of those 2 at- tacks and that we propose a unique and powerful privacy definition referred to as -diversity. additionally to assembling a proper foundation for-diversity, we have a tendency to show in AN experimental analysis that -diversity is sensible and may be enforced with efficiency.

5) Title: k-ANONYMITY: A MODEL FOR PROTECTING PRIVACY

Written By: LATANYA SWEENEY

Consider a knowledge holder, like a hospital or a bank, that encompasses a in camera control assortment of person-specific, field structured knowledge. Suppose the information holder desires to share a version of the information with researchers. However will a knowledge holder unharness a version of its personal knowledge with scientific guarantees that the people United Nations agency square measure the themes of the data can not be re-identified whereas the information stay much useful? the answer provided during this paper includes a proper protection model named face-anonymity and a collection of related policies for preparation. A unharness provides fc-anonymity protection if info for every person contained within the unharness can not be distinguished from a minimum of k-\ people whose information conjointly seems within the unharness. This paper conjointly examines re-identification attacks which will be

accomplished on releases that adhere to k- namelessness unless related policies square measure revered. The fc-anonymity protection model is vital as a result of it forms the premise on that the real-world systems called Datafly, i-Argus and fc-Similar offer guarantees of privacy protection

III.PROPOSED SYSTEM

- 1) We propose an efficient and fine-grained big data access control scheme with privacy-preserving policy, where the whole attributes are hidden in the access policy rather than only the values of the attributes.
- 2) We also design a novel Attribute Bloom Filter to evaluate whether an attribute is in the access policy and locate the exact position in the access policy if it is in the access policy.
- 3) We further give the security proof and performance evaluation of our proposed scheme, which demonstrate that our scheme can preserve the privacy from any LSSS access policy without employing much overhead.

III.SYSTEM ARCHITECTURE

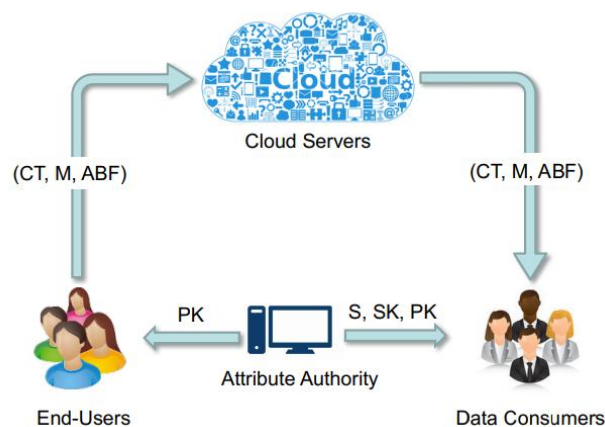


Fig.1: System Architecture

IV. GOALS AND OBJECTIVE

- The main goal of the project is to study, design and implement performance optimizations for big data frameworks. This work contributes methods and techniques to build tools for easy and efficient processing of very large data sets. It describes ways to make systems faster, by inventing ways to shorten job completion times.
- Another major goal is to facilitate the application development in distributed data-intensive computation platforms and make big-data analytics accessible to non-experts, so that users with limited programming experience can benefit from analyzing enormous datasets.
- To generate faster results.
- It reduces the complexity of data access and retrieval. When we have to dealing with big data.
- The alternative to this is apache Hadoop, which deals with big data with efficiency.
- Hadoop itself consists of Map Reduce and HDFS.
- It runs on Hadoop cluster.

V.CONCLUSION AND FUTURE SCOPE

In this paper, we've got planned associate economical and fine-grained information access management theme for large information, wherever the access policy won't leak any privacy data. completely different from the present strategies that solely part hide the attribute values within the access policies, our technique will hide the complete attribute (rather than solely its values) within the access policies. However, this could result in nice challenges and difficulties for legal information shoppers to decode information. To address this downside, we've got additionally designed associate

attribute localization formula to judge whether or not attribute is within the access policy. so as to enhance the potency, a unique Attribute Bloom Filter has been designed to find the precise row numbers of attributes within the access matrix. we've got additionally incontestible that our theme is by selection secure against chosen plaintext attacks. Moreover, we've got enforced the ABF by victimisation MurmurHash and also the access management theme to indicate that our theme will preserve the privacy from any LSSS access policy while not using abundant overhead. In our future work, we'll specialise in the way to affect the offline attribute idea attack that check the idea "attribute strings" by regularly querying the ABF

ACKNOWLEDGMENT

We might want to thank the analysts and also distributors for making their assets accessible. We additionally appreciative to commentator for their significant recommendations furthermore thank the school powers for giving the obliged base and backing.

REFERENCES

- [1] B. C. M. FUNG, K. WANG, R. CHEN, AND P. S. YU, "PRIVACY-PRESERVING DATA PUBLISHING: A SURVEY OF RECENT DEVELOPMENTS," *ACM COMPUT. SURV.*, VOL. 42, NO. 4, PP. 14:1–14:53, 2010.
- [2] X. MA, H. Li, J. Ma, Q. Jiang, S. Gao, N. Xi, and D. Lu, "Applet: A privacy-preserving framework for location-aware recommender system," *Sci China Inf Sci*, vol. 59, no. 2, pp. 1–15, 2016.
- [3] W. Xia, R. Heatherly, X. Ding, J. Li, and B. Malin, "Efficient discovery of de-identification policies through a risk-utility frontier," in *CODASPY*, 2013, pp. 59–70.
- [4] K. Benitez, G. Loukides, and B. Malin, "Beyond safe harbor: Automatic discovery of health information de-identification policy alternatives," in *IHI*, 2010, pp. 163–172.
- [5] K. E. Emam, "Heuristics for de-identifying health data," *IEEE Security and Privacy*, vol. 6, no. 4, pp. 58–61, 2008.
- [6] P. Mell and T. Grance, "The NIST definition of cloud computing," [Recommendations of the National Institute of Standards and Technology Special Publication 800-145], 2011.
- [7] R. Lu, H. Zhu, X. Liu, J. K. Liu, and J. Shao, "Toward efficient and privacy-preserving computing in big data era," *IEEE Network*, vol. 28, no. 4, pp. 46–50, 2014.
- [8] K. Yang and X. Jia, "Expressive, efficient, and revocable data access control for multi-authority cloud storage," *IEEE Trans. Parallel Distrib. Syst.*, vol. 25, no. 7, pp. 1735–1744, July 2014.
- [9] H. Li, D. Liu, K. Alharbi, S. Zhang, and X. Lin, "Enabling fine-grained access control with efficient attribute revocation and policy updating in smart grid," *KSII Transactions on Internet and Information Systems (TIIS)*, vol. 9, no. 4, pp. 1404–1423, 2015.
- [10] K. Yang, Z. Liu, X. Jia, and X. S. Shen, "Time-domain attribute-based access control for cloud-based video content sharing: A cryptographic approach," *IEEE Trans. on Multimedia* (to appear), February 2016.
- [11] B. Waters, "Ciphertext-policy attribute-based encryption: An expressive, efficient, and provably secure realization," in *Proc. of PKC'11*. Berlin, Heidelberg: Springer-Verlag, 2011, pp. 53–70.
- [12] H. Lin, Z. Cao, X. Liang, and J. Shao, "Secure threshold multi authority attribute based encryption without a central authority," in *Proc. of INDOCRYPT'08*. Springer, 2008, pp. 426–436.
- [13] T. Nishide, K. Yoneyama, and K. Ohta, "Attribute-based encryption with partially hidden encryptor-specified access structures," in *Applied cryptography and network security*. Springer, 2008, pp. 111–129.
- [14] J. Li, K. Ren, B. Zhu, and Z. Wan, "Privacy-aware attribute-based encryption with user accountability," in *Information Security*. Springer, 2009, pp. 347–362.
- [15] D. Boneh and B. Waters, "Conjunctive, subset, and range queries on encrypted data," in *Theory of cryptography*. Springer, 2007, pp. 535–554.
- [16] J. Katz, A. Sahai, and B. Waters, "Predicate encryption supporting disjunctions, polynomial equations, and inner products," in *Advances in Cryptology—EUROCRYPT'08*. Springer, 2008, pp. 146–162.
- [17] J. Lai, R. H. Deng, and Y. Li, "Fully secure ciphertext-policy hiding cpabe," in *Information Security Practice and Experience*. Springer, 2011, pp. 24–39.
- [18] L. Lei, Z. Zhong, K. Zheng, J. Chen, and H. Meng, "Challenges on wireless heterogeneous networks for mobile cloud computing," *IEEE Wireless Communications*, vol. 20, no. 3, pp. 34–44, 2013.
- [19] K. Zheng, Z. Yang, K. Zhang, P. Chatzimisios, K. Yang, and W. Xiang, "Big data-driven optimization for mobile networks toward 5g," *IEEE Network*, vol. 30, no. 1, pp. 44–51, 2016.
- [20] Z. Su, Q. Xu, and Q. Qi, "Big data in mobile social networks: a qoe-oriented framework," *IEEE Network*, vol. 30, no. 1, pp. 52–57, 2016.
- [21] H. Li, D. Liu, Y. Dai, and T. H. Luan, "Engineering searchable encryption of mobile cloud networks: when qoe meets qop," *IEEE Wireless Communications*, vol. 22, no. 4, pp. 74–80, 2015.

- [22] H. Li, Y. Yang, T. Luan, X. Liang, L. Zhou, and X. Shen, "Enabling fine-grained multi-keyword search supporting classified sub-dictionaries over encrypted cloud data," *IEEE Trans. on Dependable and Secure Computing* [DOI: 10.1109/TDSC.2015.2406704], 2015.
- [23] K. Frikken, M. Atallah, and J. Li, "Attribute-based access control with hidden policies and hidden credentials," *IEEE Trans. on Computers*, vol. 55, no. 10, pp. 1259–1270, 2006.
- [24] S. Yu, K. Ren, and W. Lou, "Attribute-based content distribution with hidden policy," in *Secure Network Protocols (NPSec'08 Workshop)*. IEEE, 2008, pp. 39–44.
- [25] J. Lai, R. H. Deng, and Y. Li, "Expressive cp-abe with partially hidden access structures," in *Proc. of ASIACCS'12*. ACM, 2012, pp. 18–19.
- [26] J. Hur, "Attribute-based secure data sharing with hidden policies in smart grid," *IEEE Trans. Parallel Distrib. Syst.*, vol. 24, no. 11, pp. 2171–2180, 2013.
- [27] A. Beimel, "Secure schemes for secret sharing and key distribution," Ph.D. dissertation, Israel Institute of Technology, Technion, Haifa, Israel, 1996.
- [28] B. H. Bloom, "Space/time trade-offs in hash coding with allowable errors," *Communications of the ACM*, vol. 13, no. 7, pp. 422–426, 1970.
- [29] K. Yang, X. Jia, and K. Ren, "Secure and verifiable policy update outsourcing for big data access control in the cloud," *IEEE Trans. Parallel Distrib. Syst.*, vol. 26, no. 12, pp. 3461–3470, Dec 2015.
- [30] C. Dong, L. Chen, and Z. Wen, "When private set intersection meets big data: an efficient and scalable protocol," in *Proc. of CCS'13*. ACM, 2013, pp. 789–800.