

Implementation of Raptor Code on GPU

Linjia Hu¹ Todor Mladenov² Saeid Noshabadi¹

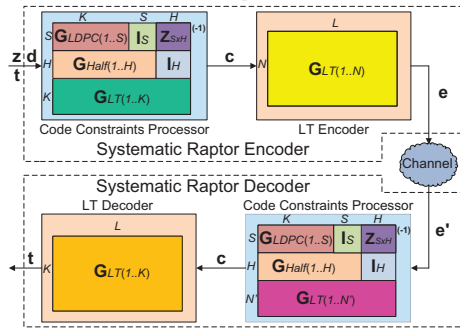
¹Department of Computer Science, Michigan Technological University, MI, USA

² Department of Information and Communication and Department of Nanobio Materials and Electronics, Gwangju Institute of Science and Technology, South Korea

Introduction

Raptor Code principle

Raptor Code come as an improvement to LT-Code, which performs as close as possible to the Shannon's channel limit and provides linear encoding and decoding time. It has been chosen for the forward error correction (FEC) scheme in 3GPP and DVB-H standards. The blow block diagram shows the systematic Raptor encoder and decoder.



At the encoding side, d denotes the input vector of the Raptor encoder, and contains K source symbols and $(S + H)$ zero symbols.

$$d_{[0:L-1]} = [z^T \ t^T]^T, \text{ where } L = K + S + H$$

Code Constraints Processor multiplies d with the inverse of the preprocessor matrix A to produce the intermediate symbols c .

$$c_{[0:L-1]} = A^{-1} \cdot d_{[0:L-1]}$$

LT Encoder can generate any number of encoded symbols e , such that

$$G_{LT} \cdot c = e_{[0:N-1]}$$

At the decoding side, we exchange the positions of Code Constraints Processor and LT Encoder.

$$c_{[0:L-1]} = [z^T \ e'^T] \cdot A^{-1} \text{ and } t_{[0:K-1]} = G_{LT} \cdot c_{[0:L-1]}$$

Matrix Inversion on CPU

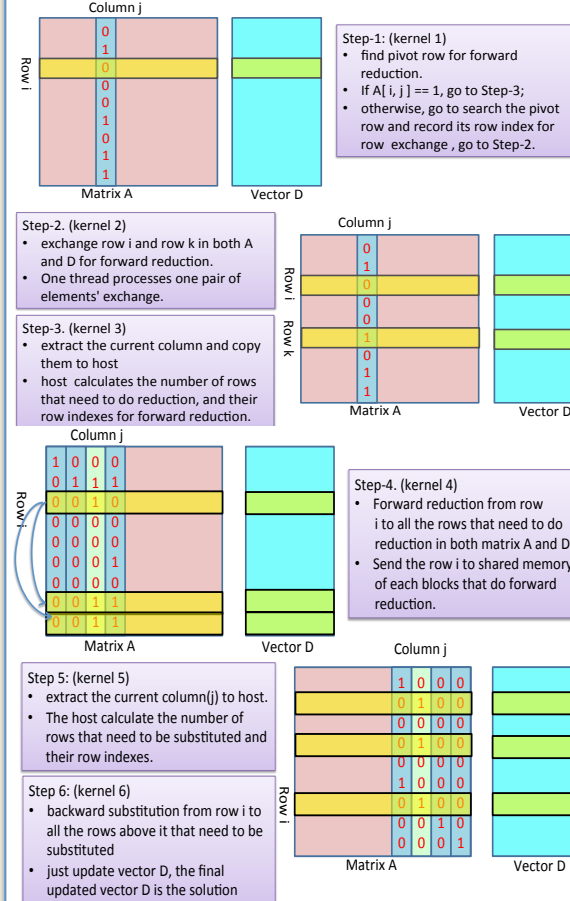
The code profiling of Raptor Code shows that the inversion of the preprocessor matrix A contributes more than 90% of the decoding time, so we concentrate on the optimization of matrix inversion algorithm. The most common matrix inversion algorithm is Gaussian Elimination (GE), so we implement GE in Galois Field $GF(2)$ on CPU as follows:

$$\begin{pmatrix} 1 & 1 & 1 & 0 \\ 1 & 1 & 0 & 1 \\ 0 & 1 & 1 & 1 \\ 1 & 1 & 1 & 1 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & 1 & 1 & 0 \\ 0 & 0 & 1 & 1 \\ 0 & 1 & 1 & 1 \\ 0 & 0 & 0 & 1 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & 1 & 1 & 0 \\ 0 & 1 & 1 & 1 \\ 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 1 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & 0 & 0 & 1 \\ 0 & 1 & 1 & 1 \\ 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 1 \end{pmatrix} \rightarrow \begin{pmatrix} 1 & 0 & 0 & 1 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 1 \end{pmatrix}$$

Implementation

Matrix Inversion on GPU

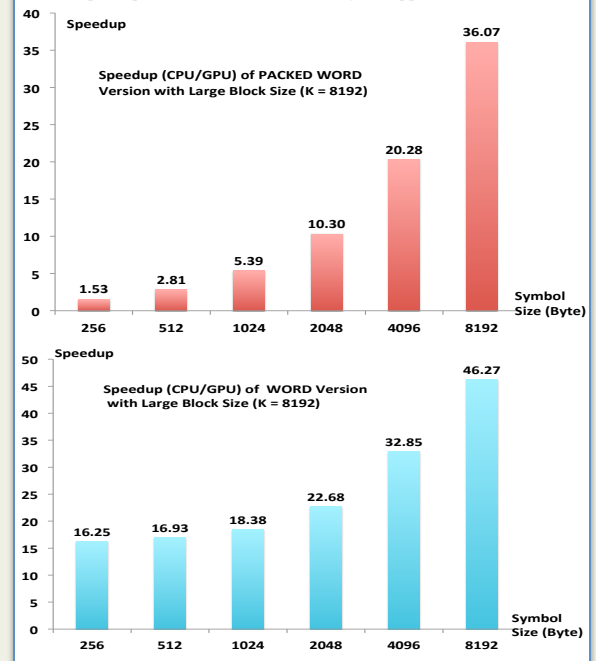
We try to implement Raptor Codes on GPU for the purpose of processing large block and symbol size efficiently. Recommended by the 3GPP and DVB-H standard, the maximum block size is 8192, and maximum symbols size is 8192 bytes. For the large matrix, we use two types of data to memory mapping. One is "WORD," which uses 32-bit word to store 1 bit matrix element, and the other is "PACKED WORD," in which 32 matrix elements are packed together into a single 32-bit word.



Conclusion

Performance

Our test platform uses a 3.20 GHz Intel Core i7 quad-core CPU, a GeForce GTX 570 graphic card with 2.5 GB video memory, CUDA 4.0 and the Fedora 13 operating system. For large block size and symbol size, the speedup of PACKED WORD version decoding can approach 36, and the speedup of WORD version decoding can approach 46.



Reference

- [1]T. Mladenov, S. Nooshabadi, and K. Kim, "Implementation and Evaluation of Raptor Codes on Embedded Systems," IEEE Transactions on Computers, Dec. 2011.
- [2]T. Mladenov, S. Nooshabadi, and K. Kim, "Strategies for the Design of Raptor Decoding in Broadcast/Multicast Delivery Systems," IEEE Transactions on Consumer Electronics, vol. 56, no. 2, pp. 423-428, May 2010