

Scenario Selection and Prediction for DVS-Aware Scheduling of Multimedia Applications

S. V. Gheorghita · T. Basten · H. Corporaal

Received: date / Accepted: date

Abstract Modern multimedia applications usually have real-time constraints and they are implemented using application-domain specific embedded processors. Dimensioning a system requires accurate estimations of resources needed by the applications. Overestimation leads to overdimensioning. For a good resource estimation, all the cases in which an application can run must be considered. To avoid an explosion in the number of different cases, those that are similar with respect to required resources are combined into, so called, *application scenarios*. This paper presents a methodology and a tool that can automatically detect the most important variables from an application and use them to select and dynamically predict scenarios, with respect to the necessary time budget, for soft real-time multimedia applications. The tool was tested for three multimedia applications. Using a proactive scenario-based dynamic voltage scheduler based on the scenarios and the runtime predictor generated by our tool, the energy consumption decreases with up to 19%, while guaranteeing a frame deadline miss ratio close to zero.

Keywords Dynamic Voltage Scheduling · Soft Real-Time · Application Scenarios · Embedded Systems

1 Introduction

Embedded systems usually contain processors that execute domain-specific programs. Many of their functionalities are

This work was supported by the Dutch Science Foundation, NWO, project FAME, number 612.064.101. More information can be found on <http://www.es.ele.tue.nl/scenarios>.

S. V. Gheorghita · T. Basten · H. Corporaal
EE Department, Electronic Systems Group
Eindhoven University of Technology
PO Box 513, 5600 MB, Eindhoven, The Netherlands
E-mail: {s.v.gheorghita,a.a.basten,h.corporaal}@tue.nl

implemented in software, which is running on one or multiple processors, leaving only the high performance functions implemented in hardware. Typical examples of embedded systems include TV sets, cellular phones and printers. The predominant workload on most of these systems is generated by stream processing applications, like video and audio decoders. Because many of these systems are real-time portable embedded systems, they have strong non-functional requirements regarding size, performance and power consumption. The requirements may be expressed as: *the cheapest, smallest and most power efficient system that can deliver the required performance*. During the design of these systems, accurate estimations of the resources needed by the application to run are required. Examples of resources include the number of execution cycles, memory-usage, and communication between application components.

Typical multimedia applications exhibit a high degree of data-dependent variability in their execution requirements. For example, the ratio of the worst case load versus the average load on a processor can be easily as high as a factor of 10 [27]. In order to save energy and still meet the real-time constraints of multimedia applications, many power-aware techniques based on dynamic voltage scaling (DVS) and dynamic power management (DPM) exploit this variability [17]. They scale the supply voltage and frequency of the processors at runtime to match the changing workload. Taking into account that the processor energy consumption depends quadratically on the supply voltage ($E \propto V_{DD}^2$), whereas its execution speed (frequency) depends linearly on the supply voltage ($f_{CLK} \propto V_{DD}$), by reducing the processor speed to half, the energy consumption can be reduced to around a quarter.

Two main broad classes of voltage and frequency scaling techniques have been developed: (i) *reactive techniques*: after a part of the application is executed, the number of

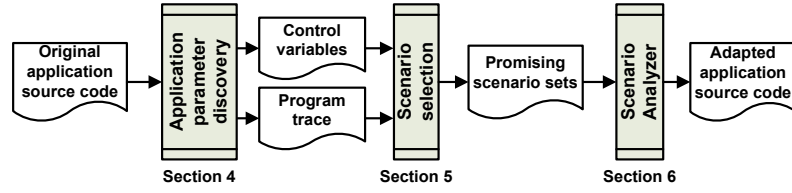


Fig. 1 Tool-flow overview.

unused processor cycles¹ is detected and the processor frequency/voltage is reduced to take advantage of the unused computation power and (ii) *proactive techniques*: detect or predict in advance that there will be unused cycles and set the processor frequency/voltage adequately. The proactive approaches are more efficient than the reactive ones, but they need a-priori derived knowledge about the input bitstream and/or the application behavior. This information can be included into the application itself as a *future case predictor* together with statically derived execution bounds for specific cases [29, 10], or it may be encoded like meta-data into the input bitstream during an offline analysis [2, 26]. To avoid an explosion in the number of different cases that are considered and in the amount of information inserted into the application or bitstream, not all different workloads are treated separately. Those that are similar with respect to required execution cycles are combined together into, so called, *application scenarios*.

Usually, to define scenarios for an application, its parameters (i.e., variables that appear in the source code) with the highest influence on the application workload are used. To the best of our knowledge, there is no way of automatically detecting these parameters, except for our previous work presented in [13, 12]. In this paper:

- we describe a method and a tool that can automatically identify the most important scenario parameters and use them to define and dynamically predict scenarios for a single-task soft real-time multimedia applications. When applied to three real-life benchmarks, the tool-flow identifies parameter sets that are similar to manually selected sets.
- we show how the method can be applied in a proactive DVS-aware scheduler, which, when applied to the three mentioned benchmarks, yields energy reductions up to 19%.

This method extends our previous work, overcoming the limitations of the static analysis for hard real-time systems used in [13], but it can not be applied to hard real-time applications, as the scenario detection and prediction is not always conservative. An earlier version of the current paper appeared as [12]. Compared to this, the current paper

¹ The unused processor cycles represent the difference between how many cycles were estimated and how many were really needed by the application.

explains the automation of scenario selection, which was a manual step in [12], at the same time slightly generalizing the scenario concept. Moreover, to overcome the fact that our approach is not conservative, we describe a runtime mechanism that guarantees the application quality, as given by the percentage of deadline misses. Finally, we evaluate our tool-flow on a larger set of benchmarks and we quantify the amount of energy that was saved by using our approach in a proactive DVS-aware scheduler.

The paper is organized as follows. Section 2 surveys related work on scenarios and different power-aware approaches for saving energy for real-time systems, and presents how our current work is different. Section 3 presents how our approach fits in a general scenario based design methodology and the kind of multimedia applications that it can be applied to. Sections 4-6 describe the three main steps of our approach of which an overview is given in figure 1. Section 7 presents the runtime calibration mechanism that is used for controlling the quality of the resulting application. In section 8, our scenario detection and prediction method is evaluated on three realistic multimedia decoders. Conclusions and future research are discussed in section 9.

2 Related work

Scenarios have been in use for a long time in different design approaches [4], including both hardware [24] and software design [9] for embedded systems. In these cases, scenarios concretely describe, in an early phase of the development process, the use of a future system. Moreover, they appear like narrative descriptions of envisioned usage episodes, or like unified modeling language (UML) use-case diagrams that enumerate, from a functional and timing point of view, all possible user actions and system reactions that are required to meet a proposed system functionality. These scenarios are called *use-case scenarios*, and characterize the system from the *user perspective*. In this work, we concentrate on a different kind of scenarios, so-called *application scenarios*, that characterize the system from the *resource usage perspective*.

The application scenario concept was first used in [33] to capture the data-dependent behavior inside a thread, to better schedule a multi-threaded application on a heterogenous multi-processor architecture, allowing the change of volt-

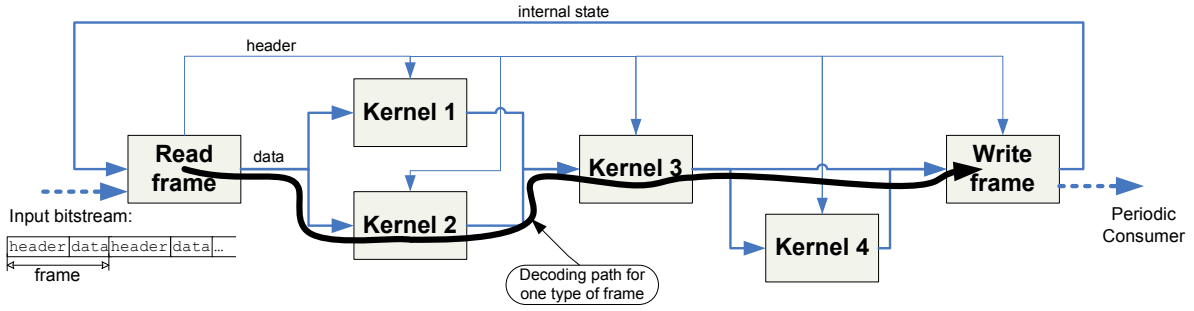


Fig. 2 Typical multimedia application decoding a frame.

age level for each individual processor. Other approaches that consider application scenarios to optimize a design include [28], [23] and [22]. In [28], the authors concentrate on saving energy for a single task application. For each manually identified scenario, they select the most energy efficient architecture configuration that can be used to meet the timing constraints. The architecture has a single processor with reconfigurable components (e.g., number and type of function units), and its supply voltage can be changed. It is not clear how scenarios are predicted at runtime. To reduce the number of memory accesses, in [23], the authors selectively duplicate parts of application source code, enabling global loop transformations across data dependent conditions. They have a systematic way of detecting the most important application behaviors based on profiling and of clustering them into scenarios based on a trade-off between the number of memory accesses and the code size increase. The final application implementation, including scenarios and the predictor, is done manually. In [22], each scenario is characterized by different communication requirements (e.g., bandwidth, latency) and traffic patterns. The paper presents a method to map a multi-task application communication to a network on chip architecture, satisfying the design constraints of each individual scenario. Most of the mentioned papers (except [23]) emphasize on how the scenarios are exploited for obtaining a more optimized design and do not go into detail on how to select and predict scenarios. Our work focuses on these last two problems.

In the context of energy saving based on DVS/DPM techniques, two different approaches exist: reactive and proactive. The proactive approaches are more efficient than the reactive ones, as they can make decisions in advance based on the knowledge about the future behavior. In order to have this knowledge available at the right moment in time, several approaches propose to a-priori process the input bitstream of a multimedia application and add to it meta-information that estimates the amount of resources needed at runtime to decode each stream object (e.g., a frame). This information is used to reconfigure the system (e.g., using DVS) in order to reduce the energy consumption, while still meeting the deadlines. In [2, 15, 26] the authors propose a

platform-dependent annotation of the bitstream, during the encoding or before uploading it from a PC to a mobile system. As it is too time expensive to use a cycle-accurate simulator to estimate the time budget necessary to decode each stream object, the presented approaches use a mathematical model to derive how many cycles are needed to decode each stream object. All these works aim at a specific application, with a specific implementation, and require that each frame header contains a few parameters that characterize the computation complexity. None of them presents a way of detecting these parameters, all assuming that the designer will provide them.

The other class of proactive approaches inserts into the application a workload case detector together with statically derived execution bounds for specific cases. The first approach for hard real-time systems was presented in [29]. It tries to predict in advance the future unused cycles, using the combined data and control flow information of the program. Its main disadvantage is the runtime overhead (which sometimes is big) that can not be controlled. In [10], we proposed a way to control this overhead, by using scenarios. We automatically detect the parameters with the highest influence on the worst case execution cycles (WCEC), and they are used to define scenarios. The static analysis used in [10] is not very powerful, as it works for some specific cases only. It is also not really suitable for soft real-time systems, as the difference between the estimated WCEC and the real number of execution cycles may be quite substantial due to the unpredictability of hardware and WCEC analysis limitations. To overcome this issue, in [12], a profiling driven approach is used to detect and characterize scenarios. It solves the issue of manually detecting parameters in the soft real-time frame-based dynamic voltage scaling algorithms, like the one presented in [25]. In this paper, we extend the approach from [12] by making the tool-flow fully automatic and more robust, and by introducing into the resulting application a runtime mechanism that controls the application quality by keeping the number of deadline misses under a required bound. Moreover, instead of only quantifying the amount of cycle-budget over-estimation reduction, we look at energy saving for a larger set of benchmarks.

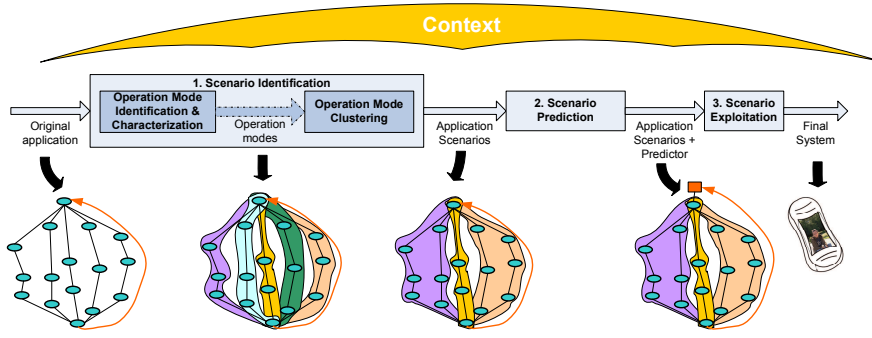


Fig. 3 Application scenario usage methodology [11].

3 Overview of our approach

This section starts by describing the characteristics of multimedia applications considered by our approach, and then details how our approach fits in the scenario methodology described in [11].

3.1 Multimedia applications

Many multimedia applications are implemented as a main loop that reads, processes and writes out individual stream objects (see figure 2). A stream object might be a bit belonging to a compressed bitstream representing a coded video clip, a macro-block, a video frame, or an audio sample. For the sake of simplicity, and without loss of generality, from now on we use the word *frame* to refer to a stream object.

The read part of the application takes the frame from the input stream and separates it into a *header* and the frame's *data*. The process part consists of several kernels. For the processing of each frame, some of these kernels are used, depending on the frame type. The write part sends the processed data to the output devices, like a screen or speakers, and saves the internal state of the application for further usage (e.g., in a video decoder, the previous decoded frame may be necessary to decode the current frame). The dynamism existing in these applications leads to the usage of different kernels for each frame, depending on the frame type. The actions executed in a particular loop iteration form an internal *operation mode* of the application. Moreover, these applications have to deliver a given throughput (number of frames per second), which imposes a time constraint (deadline) for each loop iteration. In case of soft real-time applications, a given percentage of deadline misses is acceptable.

3.2 Scenario-aware energy reduction

The scenario methodology described in [11] consists of three main steps, presented in figure 3, each of them answering to a specific question:

1. *Identification*: given an application, how is it classified into scenarios?
2. *Prediction*: given an operation mode, to which scenario does it belong?
3. *Exploitation*: given a particular scenario, what can be done to optimize the application cost in term of resource usage?

Our approach follows this methodology, in the context of saving energy using a coarse grain frame-based DVS-aware scheduling technique for soft real-time applications. In the first part of the *identification* step (*Operation mode identification and characterization*, section 4) the common operation modes are identified and profiled. As we are interested in reducing energy by exploiting the different amounts of required computation cycles of different operation modes, we identify the application variables of which the values influence the application execution time the most, and we use them to characterize the operation modes. As the number of the operation modes depends exponentially on the number of control instructions in the application, the second part of the *identification* step (*Operation mode clustering*, section 5) aims to cluster the modes into application scenarios. The described clustering algorithm takes into account factors like the cost of runtime switching between scenarios, and the fact that the amount of computation cycles for the operation modes within a scenario should always be fairly similar.

In the *scenario prediction* step (section 6) a proactive predictor is derived. Based on the parameters used to characterize the operation modes, it predicts at runtime in which scenario the application currently runs. As we aim to reduce the average energy consumption, in the *scenario exploitation* step, for each scenario, we compute the minimum processor frequency at which it can execute without missing the application's timing constraints. At runtime, when the predictor selects a new scenario, the processor frequency and supply voltage is adapted adequately. It leads to a coarse-grain schedule, as the processor frequency (and voltage) is changed once per scenario occurrence.

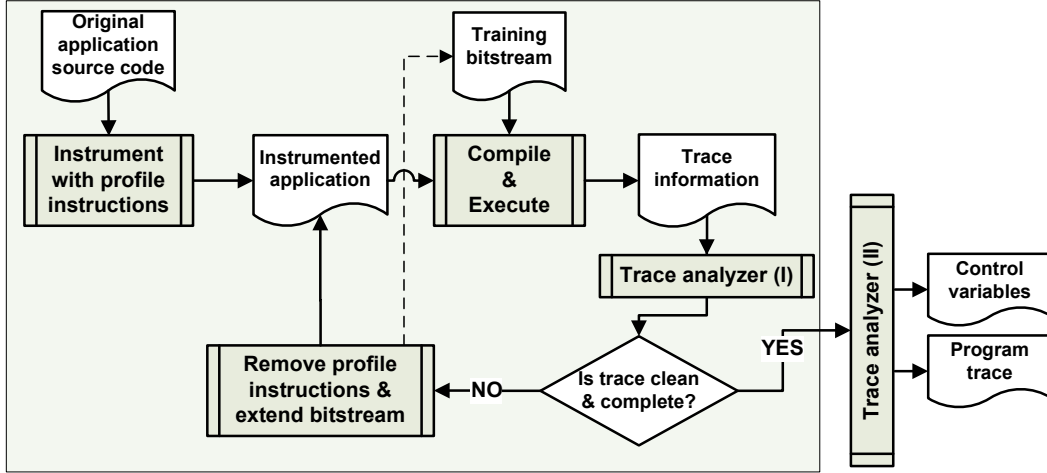


Fig. 4 Tool-flow details for deriving application parameters.

All the mentioned steps are based on profiling collected information, with the well-known limitation that the profiled information might not cover all operation modes that might occur. To overcome this limitation, a quality preservation mechanism is added to the final implementation of the application (section 7). Its role is to keep the number of deadline misses under a required threshold.

4 Application parameter discovery

This section describes the first step of our method. The method is visualized in figure 1. As explained in the previous section, it is a concrete instance of the first two steps of the methodology shown in figure 3, where the application parameter discovery step corresponds to the first part of step 1 in figure 3. This section first explains how application parameters could be used to estimate the necessary cycle budget. The remaining parts of the section detail how these parameters are discovered by our method.

4.1 Cycle budget estimation

During system design, accurate estimations of the resources needed by the application in order to meet the desired throughput are required. This paper focuses on the cycle budget needed to decode a frame in a specific period of time (p_{frame}) on a given single-processor platform. This budget depends on the frame itself and the internal state of the application. In relevant related work [2,15,26], it is typically assumed that the cycle budget $c(i)$ for frame i can be estimated using a linear function on data-dependent arguments with data-independent, possibly platform dependent, coefficients:

$$c(i) = C_0 + \sum_{k=1}^n C_k \xi_k(i), \quad (1)$$

where the C_k are constant coefficients that usually depend on the processor type, and the $\xi_k(i)$ are n arguments that depend on the frame i from the input bitstream². Using for each frame its own transformation function with all possible source-code variables as data-dependent arguments, gives the most accurate estimates. However, this approach leads to a huge number of very large functions. To reduce the explosion in the number of functions, the frames with small variation in decoding cycles are treated together, being combined in *application scenarios*. To reduce the size of each function, only the variables whose values have a large influence on the decoding time of a frame should be used. The following subsections present a method to identify these variables.

4.2 Control variable identification

The variables that appear in an application may be divided into *control variables* and *data variables*. Based on the control variable values, different paths of the application are executed, as they determine, for example, which conditional branch is taken or how many times a loop will iterate. The data variables represent the data processed by the application. Usually, the data variables appear as elements of large arrays, implicitly or explicitly declared. Attached to each array, there can be a control variable that represents the array size. Considering that each element of a data array is one data variable, it can be easily observed that, usually, there are a lot more data variables than control variables in a multimedia application.

The control variables are the ones that influence the execution time of the program the most, as they decide how

² Equation 1 could potentially have non-linear dependencies on the $\xi_k(i)$ (e.g., $\xi_k(i)^2$). For this paper, the function format is not relevant, as we only use the $\xi_k(i)$ to predict the program scenarios and not to estimate the cycle count.


```

void process(char *a, int n) {
1   int i = 0;
2   while(i < n) {
3       f(a[i]);
4       f(a[a[i]]);
5       i++;
6   }
7 }

```

Fig. 5 An educational example.

often each part of the program is executed. Therefore, as our scope is to identify a small set of variables that can be used to estimate the amount of cycles required to process a frame, we separate the variables into data and control, based on application profiling. Moreover, we identify a subset of the control variables that do not influence the execution time and hence are not of interest to us. Both aspects are handled by the trace analyzer discussed in the next subsection.

The large gray box in figure 4 shows the work-flow for control variable identification. It starts from the application source code which is then instrumented with profile instructions for all read and write operations on the variables. The instrumented code is compiled and executed on a training bitstream and the resulting program trace is collected and analyzed. To find a *representative* training bitstream that covers most of the behaviors which may appear during the application life-time, particularly including the most frequent ones, is in general a difficult problem. However, an approach similar to the one presented in [19], where the authors show a technique for classifying different multimedia streams, could be used. The analysis performed on the collected trace information aims to discover if the trace contains data variables. If any are discovered, the profile instructions that generate this information are removed from the source code, and the process of compiling, executing and analyzing is repeated until the trace does not contain data variables anymore. As our method generates a huge trace if it is applied from the beginning on a large bitstream, we start with a few frames of the bitstream in the first iteration. At each iteration, we increase the number of considered frames as the size of trace information generated per frame reduces. The process is complete if the entire training bitstream is processed and the resulting trace does not contain any data variables.

4.3 Trace analyzer

The trace analyzer has two roles: (i) at each iteration of the flow for control variable identification, it identifies data variables and control variables that do not affect execution time substantially; and (ii) when the process is complete, it generates the data necessary for the scenario selection step explained in section 5 and a list of the remaining control variables.

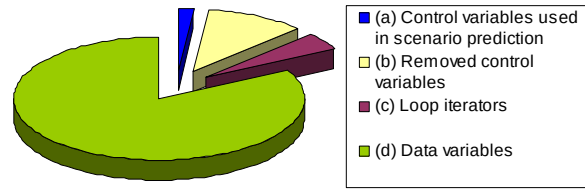


Fig. 6 Variable distribution for MP3.

The data variables that are declared as explicit arrays can be found via a straightforward static analysis of the source code. For the rest of the data variables, stored in implicitly declared arrays (e.g., the variable *a* from the source code of figure 5), the trace analyzer applies the following rule: if in the trace information generated for each frame, there is a program instruction that reads or writes a number of different memory addresses (e.g., the instructions from lines 3 and 4 in figure 5) larger than a threshold, we consider that all these memory addresses are linked to data variables, as this operation looks like accessing a data array. For this decision, we do not look for a specific array access pattern (e.g., a sequential access pattern as in line 3 or a random access pattern as in line 4 of our example). The profiling in combination with a threshold allows to differentiate between implicitly declared arrays that store data or control variables. This can not be obtained only by inspecting the source code, due to the complexity of the C language and the limitation of existing static analysis techniques, like pointer alias analysis [14]. Based on practical experience, we observed that the threshold is quite low. It is a configuration parameter for our tool, and its default value is four, as it is the appropriate value found by us in practice.

Loop iterators are the control variables that we consider to have only a small influence on the application execution time and that are easy to identify based on the trace information generated for each frame. These variables are not used to decide how many times a loop iterates; they just count the number of iterations. For example, in the piece of code of figure 5, the variable *n* bounds the number of iterations, while the loop iterator *i* counts them. Variable *n* might be of interest, but *i* is not. If there is a program instruction that writes the same variable more than once, this variable can be considered a loop iterator³.

When the trace analyzer finishes, all data variables and loop iterators are removed. The trace analyzer generates a list with the remaining variables from the trace which are candidates for the ξ_k used in equation (1). During the scenario analyzer step (section 6), their number is further re-

³ The same behavior appears also in the case of counters, but we do not make the difference between counters and iterators, removing these variables in both cases.

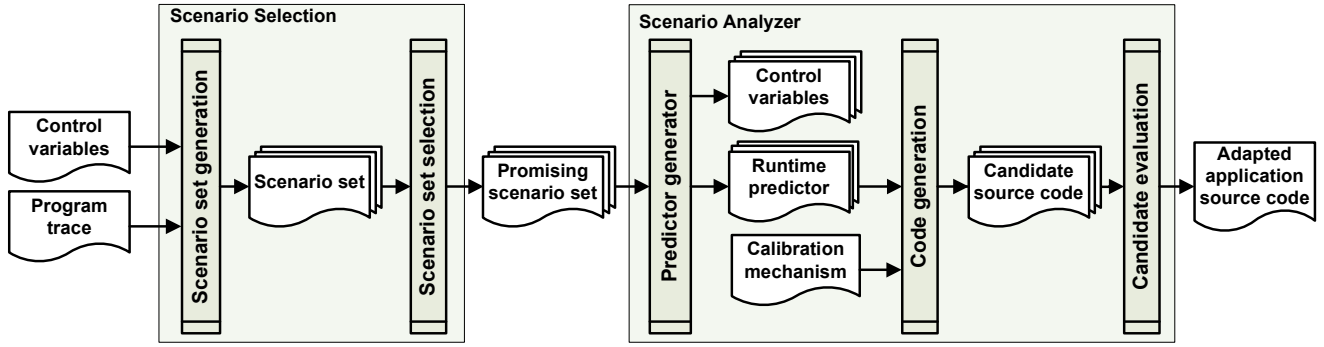


Fig. 7 Tool-flow details for scenario selection and analyzer steps.

duced. Figure 6 shows the categories into which the application variables are divided, where category (b) covers the variables removed during the scenario analyzer step.

Besides the write and read operations, the program trace contains also the number of cycles needed to decode each frame (part of the operation mode characterization in step 1 of figure 3). This information is used in the scenario selection step, discussed in the next section.

5 Scenario selection

This section presents our scenario selection approach (the second step in figure 1 and part 2 of step 1 in figure 3). It first details the scenario selection problem. It then continues in section 5.2 by introducing the frame and scenario signatures that capture all the relevant information needed for scenario selection and prediction. The remaining part of the section describes the actual scenario selection step, which is detailed in the left gray box of figure 7. It consists of two main processes: (i) using a heuristic approach, multiple scenario sets are generated from the information previously derived by profiling the training bitstream (section 5.3), and (ii) from the generated scenario sets the most promising ones from an energy saving point of view are selected (section 5.4).

5.1 The scenario selection problem

In [12], scenarios are manually identified based on a graphically depicted distribution histogram that shows on the horizontal axis the number of cycles needed to decode a frame and on the vertical axis how often this cycle budget was needed for the training bitstream. Each identified scenario j is characterized by a cycle budget interval $(c_{lb}(j), c_{ub}(j)]$ that bounds the number of cycles needed to decode each frame that is part of the scenario. The set of identified scenarios covers all the frames that appear in the training bitstream.

In the final application source code generated by our method, for each frame of a scenario, c_{ub} is used as an esti-

mate for the required cycle budget for processing it. So, each scenario introduces an over-estimation that is determined by the difference between c_{ub} and the average amount of cycles needed to process the frames belonging to it. As the aim is to exploit DVS, for each scenario the targeted processor frequency is set to the lowest frequency that can deliver at least c_{ub} cycles within a p_{frame} period of time. An overhead of maximum t_{switch} seconds⁴ is taken into account for changing the processor frequency at runtime, when the application switches between scenarios. So, tight bounds c_{ub} and limited switching frequency are important.

Manual scenario selection is a time-consuming iterative job. The process starts by deriving an initial set of scenarios from the distribution histogram. Then, its quality in prediction and over-estimation is evaluated. It might not be straightforward to unambiguously characterize the manually selected scenarios by means of the variables identified in the previous section. Based on the obtained results, the set can be adapted and re-evaluated as often as necessary. A manual selection approach, similar to the one presented in [12], can easily exploit the information that can be extracted from the distribution histogram: (i) how often scenarios occur at runtime and (ii) the introduced cycle-budget over-estimation. However, it is very difficult, even impossible, to take into account other necessary ingredients for selecting the best set of scenarios that are runtime detectable and introduce the lowest over-estimation, such as: (i) whether it is possible to distinguish at runtime between scenarios based on the considered control variables, (ii) the possible overlap in the cycle budget intervals of identified scenarios, (iii) how many switches appear between each two scenarios, and (iv) the runtime scenario prediction and system reconfiguration (i.e., voltage/frequency scaling) overhead. All this information is taken into account in the heuristic algorithm presented in the following subsections. A running example, a simplified MPEG-2 motion compensation (MC) task, is used throughout the section for easier understanding.

⁴ t_{switch} can be extracted from the processor datasheet.

$$\begin{aligned}
\Sigma_f(1) &= (V_{f(1)} = \{(\xi_1, 1), (\xi_2, \sim), (\xi_3, 2)\}, 40) & \Sigma_f(4) &= (V_{f(4)} = \{(\xi_1, 2), (\xi_2, 352), (\xi_3, 12)\}, 112) & \Sigma_f(7) &= (V_{f(7)} = \{(\xi_1, 2), (\xi_2, 704), (\xi_3, 12)\}, 108) \\
\Sigma_f(2) &= (V_{f(2)} = \{(\xi_1, 2), (\xi_2, 352), (\xi_3, 2)\}, 39) & \Sigma_f(5) &= (V_{f(5)} = \{(\xi_1, 2), (\xi_2, 352), (\xi_3, 4)\}, 42) & \Sigma_f(8) &= (V_{f(8)} = \{(\xi_1, 2), (\xi_2, 704), (\xi_3, 4)\}, 41) \\
\Sigma_f(3) &= (V_{f(3)} = \{(\xi_1, 1), (\xi_2, \sim), (\xi_3, 12)\}, 110) & \Sigma_f(6) &= (V_{f(6)} = \{(\xi_1, 2), (\xi_2, 704), (\xi_3, 2)\}, 39)
\end{aligned}$$

Fig. 8 A sequence of frame signatures

$$\begin{aligned}
& F_{j_1} = \{1, 2, 6\} \quad \Sigma_s(j_1) = ([39, 40], 2, 3, 2) & s(j_1, j_2) &= 0 & s(j_2, j_1) &= 1 & j = cls(j_1, j_2) \quad F_j = \{1, 2, 5, 6, 8\} \\
& F_{j_2} = \{5, 8\} \quad \Sigma_s(j_2) = ([41, 42], 1, 2, 2) & o(j_1, j_2) &= o(j_2, j_1) = 2 + 1 + 2 \cdot 3 = 9 & \Sigma_s(j) &= ([39, 42], 9, 5, 3) \\
& \text{(a) Signatures} & \text{(b) Functions} & \text{(c) Clustering} \\
& t_{switch} = 1\mu sec \quad p_{frame} = 10\mu sec \quad sw(j) = \lceil (42/10) \cdot 1 \rceil = 5 & & & & \\
& u_{ub}(j) = \lceil (3 \cdot 5 - 9)/5 \rceil = 2 \quad c_{ub}(j) = 42 + 2 = 44 & & & & \\
& \text{(d) Upper bound adaptation} & & & & \text{(e) Clustering cost} \\
& & & & & sw(j_1) = 4 \quad sw(j_2) = 5 \quad \alpha = 1 \\
& & & & & u_{ub}(j_1) = \lceil \frac{2 \cdot 4 - 2}{3} \rceil = 2 \quad u_{ub}(j_2) = \lceil \frac{2 \cdot 5 - 1}{2} \rceil = 5 \quad u_{ub}(j) = \lceil \frac{3 \cdot 5 - 9}{5} \rceil = 2 \\
& & & & & cost(j) = 9 - 2 - 1 - 1 \cdot (0 \cdot 4 + 1 \cdot 5) + 2 \cdot (3 + 2) - 2 \cdot 3 - 5 \cdot 2 = -5
\end{aligned}$$

Fig. 9 Example of scenarios

5.2 Scenario signatures

It is our aim to derive scenarios and scenario predictors from the knowledge that can be extracted from the training bitstream. To this end, we first characterize each frame from the training bitstream in terms of the control variables and its cycle count. This information is used in both the scenario selection and analyzer steps.

Let C be the set of control variables ξ_k obtained through the trace analyzer. **Frame signatures** are obtained by processing the trace generated for the training bitstream. For a frame i its signature $\Sigma_f(i)$ is defined as a pair:

$$\Sigma_f(i) = (V_{f(i)} = \{(\xi_k, \xi_k(i)) | \xi_k \in C\}, c(i)), \quad (2)$$

where $\xi_k(i)$ is the value of control variable ξ_k for frame i , and $c(i)$ represents the number of cycles used to process frame i . For each frame, there can be some variables ξ_k that are not accessed during its processing, so they have undefined values. An example of a sequence of frame signatures for a training bitstream is shown in figure 8, where $\sim$ represents an undefined value.

Assume, for the moment, that all frames in the training bitstream have been partitioned into a set of scenarios. Let F_j be the set of all frames that belong to scenario j . A **scenario signature** can then be computed from the signature of all the frames in the training bitstream that are part of the scenario. Scenario signatures quantify the aspects of a scenario that are used in the scenario selection. For a scenario j , its scenario signature $\Sigma_s(j)$ is defined as a 4-tuple:

$$\Sigma_s(j) = ([c_{lb}(j), c_{ub}(j)], o(j), f(j), s(j)), \quad (3)$$

where $c_{lb}(j) = \min_{i \in F_j}(c(i))$ and $c_{ub}(j) = \max_{i \in F_j}(c(i))$ bound the number of cycles needed to process each frame part of the scenario; $o(j) = \sum_{i \in F_j}(c_{ub}(j) - c(i))$ represents the accumulated cycle budget over-estimation that this scenario introduces for the training bitstream; $f(j)$ counts how often the scenario appears (i.e., $f(j)$ equals the cardinality of F_j); and $s(j)$ counts how many times the application

switches from this scenario to other scenarios (i.e., it counts in the training bitstream the number of frame intervals that consist of frames in scenario j). Figure 9(a) gives an example of two scenarios that contain some of the frames presented in figure 8.

The scenario selection algorithm repeatedly considers scenario candidates for clustering into one new scenario. To derive the signature for the scenario resulting from clustering a pair of scenarios (j_1, j_2) , we introduce:

- $s(j_1, j_2)$ is the number of times that the application switches from scenario j_1 to scenario j_2 while processing the training bitstream, with $s(j_1, j_2) = 0$ if $j_1 = j_2$;
- $o(j_1, j_2)$ is the over-estimation introduced by clustering the two scenarios into a single one, where

$$\begin{aligned}
o(j_1, j_2) &= o(j_1) + o(j_2) + \\
&\quad \begin{cases} (c_{ub}(j_1) - c_{ub}(j_2)) \cdot f(j_2), & \text{if } c_{ub}(j_1) > c_{ub}(j_2) \\ (c_{ub}(j_2) - c_{ub}(j_1)) \cdot f(j_1), & \text{if } c_{ub}(j_1) \leq c_{ub}(j_2) \end{cases}
\end{aligned} \quad (4)$$

Figure 9(b) gives a numerical example of how these functions are computed for the scenarios from figure 9(a) and the frame sequence given in figure 8.

Given two scenarios j_1 and j_2 , with signatures $\Sigma_s(j_1)$ and $\Sigma_s(j_2)$, their **clustering** is a scenario $cls(j_1, j_2)$ with the signature:

$$\begin{aligned}
\Sigma_s(cls(j_1, j_2)) &= \\
& ([\min(c_{lb}(j_1), c_{lb}(j_2)), \max(c_{ub}(j_1), c_{ub}(j_2))], \\
& o(j_1, j_2), \\
& f(j_1) + f(j_2), \\
& s(j_1) + s(j_2) - s(j_1, j_2) - s(j_2, j_1)).
\end{aligned} \quad (5)$$

Figure 9(c) displays the scenario resulting from clustering the scenarios in figure 9(a).

5.3 Scenario sets generation

This step, of which pseudo-code is shown in figure 10, represents the first part of the *scenario selection algorithm*.


```

GENERATESCENARIOSETS(VECTOR frames)
1  solutions ← ∅
2  scenarioSet ← INITIALCLUSTERING(frames)
3  solutions.INSERT(scenarioSet)
4  while (scenarioSet.SIZE() ≠ 1)
5      do (j1, j2) ← GETTWO SCENARIOS TO CLUSTER(scenarioSet)
6          j ← CLUSTERSCENARIOS(j1, j2)
7          scenarioSet.REMOVE(j1)
8          scenarioSet.REMOVE(j2)
9          scenarioSet.INSERT(j)
10         solutions.INSERT(scenarioSet)
11  for each scenarioSet in solutions
12      do for each s in scenarioSet
13          do ADAPTSCENARIOBOUNDS(s)
14  return solutions

```

Fig. 10 The scenario sets generation algorithm.

Its role is to divide the execution cases of the application in a number of scenarios. It receives as parameter the vector of frame signatures for the training bitstream. The algorithm returns multiple scenario sets, each of them covering all the given frames and being a potentially promising solution that represents a trade-off between the number of scenarios and the introduced over-estimation. More scenarios lead to less over-estimation. However, more scenarios lead to more switches and a larger predictor, which may increase the cycle overhead and enlarge the application source code too much.

In the initialization phase (line 2), the algorithm generates an initial set of scenarios. It takes into account that there is no way to differentiate at runtime between two frames i_1 and i_2 if their signatures are such that $V_{f(i_1)} = V_{f(i_2)}$. So, in the initialization phase, all the frames i that have in the signature the same set $V_f(i)$ are clustered together in the same scenario.

The processing part of the algorithm starts with the initial set of scenarios and it is repeated until the scenario set contains only one scenario that clusters together all frames. At each iteration, the two most promising scenarios to be clustered are selected using a heuristic function, discussed in more detail below, and they are replaced in the scenario set by the scenario resulting from their clustering.

After the processing part, for each scenario j from each set of scenarios (lines 11-13), the upper bound of the cycle budget interval $c_{ub}(j)$ is adapted to accommodate, on average, the cycles spent to switch from this scenario to other scenarios. The maximum number of cycles used to switch from j is given by:

$$sw(j) = \lceil (c_{ub}(j) / p_{frame}) \cdot t_{switch} \rceil, \quad (6)$$

where p_{frame} is the frame period, $c_{ub}(j) / p_{frame}$ is the processor frequency at which the scenario j is executed and t_{switch} is the maximum time overhead introduced by a frequency switching. In principle, the over-estimation introduced by a scenario can be used to accommodate for switching cycles. However, this over-estimation may be too small.

Thus, if the over-estimation $o(j)$ introduced by the scenario is smaller than the total number of processor cycles needed to switch from it to other scenarios ($s(j) \cdot sw(j)$), then $c_{ub}(j)$ is incremented. Otherwise, it remains unchanged. The following formula computes the incrementing value:

$$u_{ub}(j) = \max \left(\left\lceil \frac{s(j) \cdot sw(j) - o(j)}{f(j)} \right\rceil, 0 \right). \quad (7)$$

In figure 9(d) the cycle budget upper bound is recomputed for the scenario defined in Figure 9(c).

Recall that the aim of this work is to save energy. The tested heuristic functions for selecting which scenarios to cluster are based on cost functions that take into account: (i) the over-estimation of the resulting scenario, (ii) the cycle budget upper bound adaptation that should be done for each scenario, and (iii) the number of switches between scenarios and the switching overhead. Via the aspects (i) and (ii), it is taken into account that the over-estimation introduced by a scenario could be used to compensate for the switching overhead from this scenario to other scenarios. There is a one-to-one correspondence between cost incurred by over-estimation cycles and cycles lost or gained via budget adaptation. Switching cost (aspect iii) will generally decrease when clustering scenarios. However, switching cost given in cycles should be weighted because the energy cost of these cycles depends on the ratio between the energy consumed during the frequency switching, information that can be taken from the processor datasheet, and the amount of energy used by normal processor operation during a period of time equal to t_{switch} . Considering all these aspects, the most promising clustering heuristic function that we found selects the pair of scenarios with the lowest cost taken as *over-estimation minus weighted switching plus adaptation*. Our experiments show that this cost function gives good results, while dropping any of the three main aspects gives worse results. Formally, for scenarios j_1 and j_2 the clustering cost is given by:

$$\begin{aligned}
 cost(c_{ls}(j_1, j_2)) = & \\
 & o(j_1, j_2) - o(j_1) - o(j_2) \\
 & - \alpha \cdot (s(j_1, j_2) \cdot sw(j_1) + s(j_2, j_1) \cdot sw(j_2)) \\
 & + u_{ub}(c_{ls}(j_1, j_2)) \cdot (f(j_1) + f(j_2)) \\
 & - u_{ub}(j_1) \cdot f(j_1) - u_{ub}(j_2) \cdot f(j_2),
 \end{aligned} \quad (8)$$

where α is a weighting coefficient for the number of cycles gained by reducing the number of switches. Figure 9(e) shows how the cost is computed for the two scenarios defined in Figure 9(a).

5.4 Scenario sets selection

This second and last step of the *scenario selection algorithm* aims to reduce the number of solutions that should be further evaluated, as the evaluation of each set of scenarios is

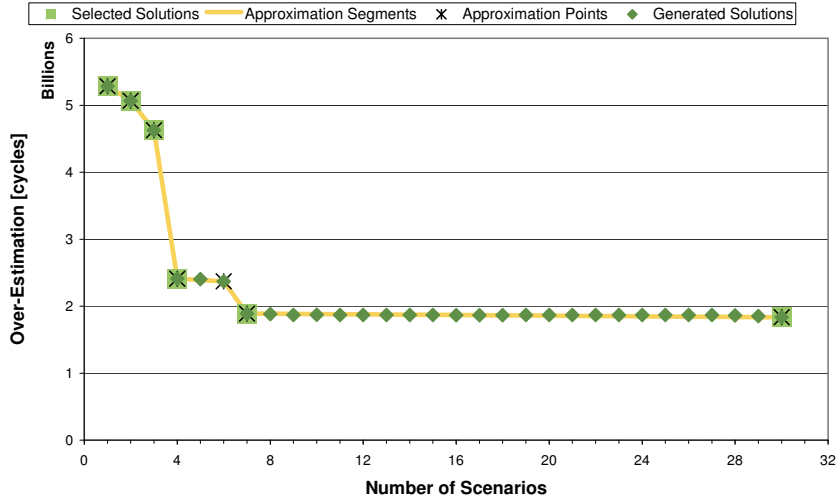


Fig. 11 Scenario sets selection for MPEG-2 MC based on over-estimation.

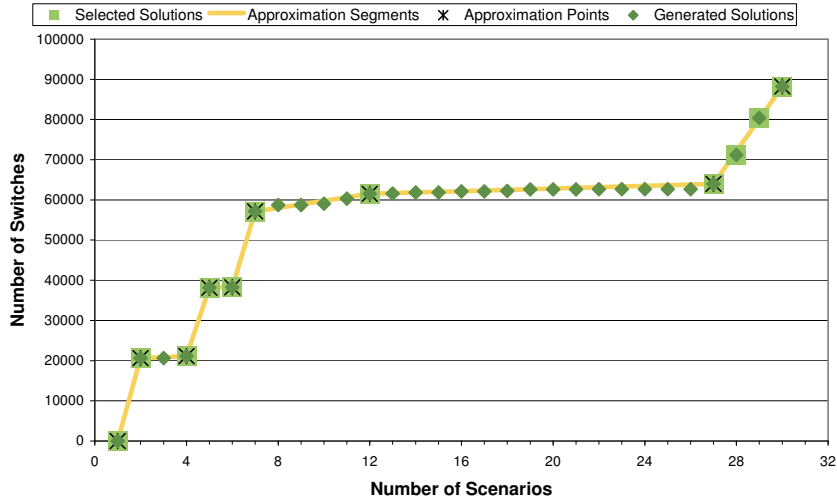


Fig. 12 Scenario sets selection for MPEG-2 MC based on number of switches.

a time-consuming operation. It chooses from the previously generated sets of scenarios the most promising ones. The goal is to find interesting trade-offs in cost (code size and runtime overhead) and gains (cycles and energy). Therefore, for making this decision, for each scenario set, the amount of introduced over-estimation and the number of runtime scenario switches are taken into account. Each solution is considered as a point in two 2-dimensional trade-off spaces: (i) the number of scenarios (m) versus introduced over-estimation ($\sum_{j=1}^m o(j)$), and (ii) the number of scenarios versus the number of runtime switches ($\sum_{j_1=1}^m \sum_{j_2=1}^m s(j_1, j_2)$). In the example given in figures 11 and 12 these points are called *generated solutions*. Each of the two charts is independently used to select a set containing promising solutions, and finally the two sets are merged. The selection algorithm consists of five steps:

1. For each chart, the sequence of solutions, sorted according to the number of scenarios, is approximated with a

set of line segments, each of them linking two points of the set, such that the sum of the squared distances from each solution to the segment used to approximate it is minimized. This problem is an instance of the *change detection* problem from the data mining and statistics fields [5]. To avoid the trivial solution of having a different segment linking each pair of consecutive points, a penalty is added for each extra used segment. In figures 11 and 12, the selected segments and their end points are called *approximation segments/points*.

2. For each chart, we initially select all the approximation points to be part of the chart's set of promising solutions. These points are potentially interesting because they correspond to solutions in the trade-off spaces where the trends in the development in over-estimation (figure 11) and number of runtime switches (figure 12) change.
3. For each approximation segment from the over-estimation chart, its slope is computed. If it is very small

compared to the slope of the entire sequence of solutions⁵, its right end point is removed from the set of promising solutions, as for similar over-estimation, we would like to have the smallest number of scenarios because that reduces code size and switches. In figure 11, for the segment between the solutions with 4 respectively 6 scenarios, the solution with 6 scenarios is discarded. The same rule does not apply for the *switches chart* because both end points are of interest. For a similar number of switches, the right end point represents the solution with the lowest over-estimation, and the left end point is the solution with the smallest predictor.

4. For each approximation segment from each chart, if its slope is larger than the slope of the entire sequence of solutions, intermediate points, if they exist, may be selected. They represent an interesting trade-off between the number of scenarios and the potential gains in over-estimation or number of switches. The percentage of selected points is chosen to depend on the ratio between the two slopes. In figure 12, the solutions with 28 and 29 scenarios are selected as intermediate points.
5. The sets of promising solutions generated for the trade-off spaces are merged, and the resulting union represents the set of the most promising solutions that will be further evaluated.

6 Scenario analyzer

The scenario analyzer step is detailed in the right gray box from figure 7. It corresponds to the third step in figure 1, and it is an instance of step 2 of the general methodology of figure 3. It starts from the previous selected set of solutions, each solution being a set of scenarios that covers the whole application. For each solution, it generates: (i) for each scenario, an equation that characterizes the scenario depending on the application control variables; (ii) the source code of the predictor that can be used to predict at runtime in which scenario the application is running; and (iii) the list of the variables used by this predictor. The predictor together with the runtime quality calibration mechanism described in section 7 is used to generate the source code for each solution. The best application implementation is selected by measuring the energy saving of each generated version of the source code on the training bitstream.

Scenario characteristic function: For each frame i , using its signature as defined in section 5.2, a boolean function $\chi_f(i)$ over variables ξ_k characterizing the frame is defined:

$$\chi_f(i)(\vec{\xi}_k) = \bigwedge_k (\xi_k = \xi_k(i)). \quad (9)$$

⁵ The sequence slope is the slope of the segment that links the first and the last point from the sequence.

By using these functions, for each scenario j , a boolean function $\chi_s(j)$ over variables ξ_k characterizing the scenario is defined. Recall that F_j denotes the set of frames belonging to scenario j .

$$\chi_s(j)(\vec{\xi}_k) = \bigvee_{i \in F_j} \chi_f(i)(\vec{\xi}_k). \quad (10)$$

The canonical form of this boolean function is obtained using the Quine McCluskey algorithm [20]. These functions can be used at runtime to check for each frame in which scenario the application should execute. Based on the *initial clustering* from the *scenario selection* step, at most one of these functions evaluates to *true* when applied to the control variable values of a frame. However, because these functions are computed based on a training bitstream, a special case may appear when a new frame i is checked against them: no scenario j for which $\chi_s(j)(\vec{\xi}_k(i))$ evaluates to *true* exists. In this case, the frame is classified to be in the so-called *backup scenario*, which is the scenario j with the largest $c_{ub}(j)$ among all the scenarios.

Runtime predictor: The operations that change the values of the variables ξ_k are identified in the source code. Using a static analysis, for each of the possible paths within the main loop of the multimedia application, the instruction that is the last one to change the value of any variable ξ_k is identified. After this instruction, the values of all required variables are known. An identical runtime predictor is inserted after each of such instructions. This leads to multiple mutually exclusive predictors, from which precisely one is executed in each main loop iteration to predict the current scenario. An extension is to consider refinement predictors active at multiple points in the code to predict the current scenario: the first one will detect a set of possible scenarios, and the following will refine the set until only one scenario remains. This extension might save more energy, as earlier switching between scenarios may be done. However, we leave this point open for future research.

We can use as the runtime predictor the scenario equations derived above. However, for a faster runtime evaluation, code optimization and the possibility of introducing more flexibility in the prediction, a decision diagram is more efficient. So, we derive the runtime predictor as a *multi-valued decision diagram* [32], defined by a function

$$f : \Omega_1 \times \Omega_2 \times \dots \times \Omega_n \rightarrow \{1, \dots, m\}, \quad (11)$$

where Ω_k is the set of all possible values of the type of variable ξ_k (including $\sim$ that represents undefined) and m is the number of scenarios in which the application was divided. The function f maps each frame i , based on the variable values $\xi_k(i)$ associated with it, to the scenario to which the frame belongs. The decision diagram consists of a directed acyclic graph $G = (V, E)$ and a labeling of the nodes and

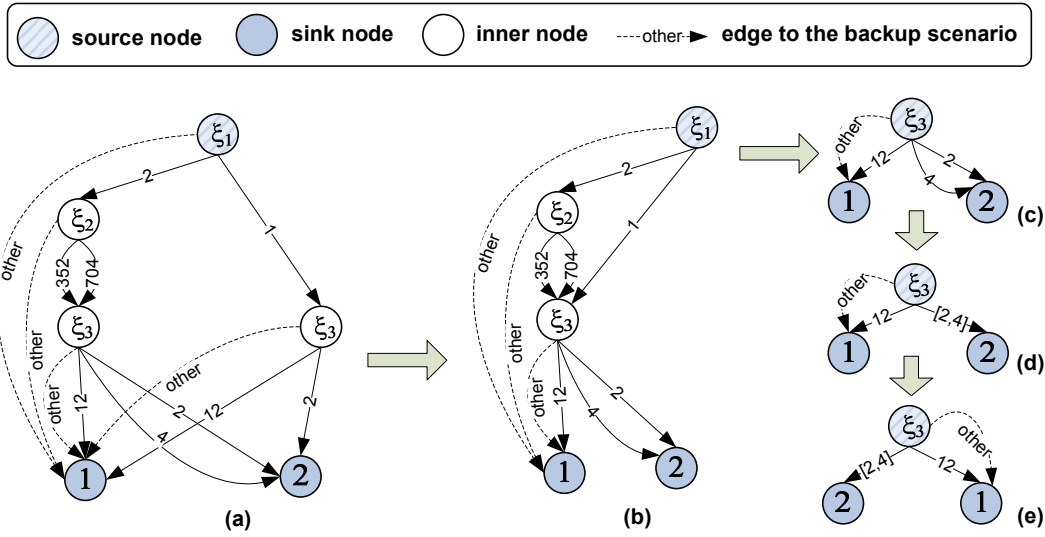


Fig. 13 Simplified MPEG-2 MC decision diagrams: (a) original; (b) merging ξ_3 ; (c) removal of ξ_1 and ξ_2 ; (d) intervals; (e) reorder.

edges. The sink nodes get labels from $1, \dots, m$ and the inner (non-sink) nodes get labels from $\xi_1, \dots, \xi_n$. Each inner node labeled with ξ_k has a number of outgoing edges equal to the number of the different values $\xi_k(i)$ that appear for variable ξ_k in all frames from the training bitstream plus an edge labeled with *other* that leads directly to the backup scenario. This edge is introduced to handle the case when, for a frame i , there is no scenario j for which $\chi_s(j)(\xi_k(i))$ evaluates to *true*. Only one inner node without incoming edges exists in V , which is the source node of the diagram, and from which the diagram evaluation always starts. On each path from the source node to a sink node each variable ξ_k occurs at most once. An example of a decision diagram for the sequence of frames of figure 8 is shown in figure 13(a).

When the decision diagram is used in the source code to predict the future scenario, it introduces two additional cost factors: (i) *decision diagram code size* and (ii) *average evaluation runtime cost*. Both can be measured in number of comparisons. To reduce the decision diagram size, a tradeoff with the decision quality is done. All the optimization steps done in our decision diagram generation algorithm (figure 14) are based on practical observations. The algorithm consists of five main steps:

1. *Initial decision diagram construction* (lines 1-21): For each scenario, a node is created and introduced in the decision diagram, and the node for the backup scenario is saved for future use (lines 2-4). For each node, the following information is stored: (i) the set of frames of the training bitstream for which the scenario prediction process passes through the node, (ii) its label (a control variable or a scenario identifier), (iii) its type (SOURCE, SINK and INNER) and (iv) the variables that were not used as labels for the nodes on the path from the source node. For SINK nodes, the latter is irrelevant, and hence these

```
NODE::NODE(SET frames, STRING label, NODETYPE type, SET vars);
```

```
GENERATEDECISIONDIAGRAM(SET frames, SET scenarios,  
                          SCENARIO backup, SET vars)
```

```
1 dd ← new DECISIONDIAGRAM()
2 for each s in scenarios
3   do dd.INSERT(new NODE(0, s.name, SINK, 0))
4   b ← dd.GETNODE(backup.name)
5 nodes ← new LIST()
6 nodes.PUSH(new NODE(frames, NIL, SOURCE, vars))
7 while (nodes.SIZE() > 0)
8   do n ← nodes.POP()
9   ξ ← n.GETVAR()
10  n.label ← ξ.name
11  _vars ← n.vars - ξ
12  for each v in ξ.values
13    do _frames ← n.frames.GETFRAMES(ξ = v)
14    if (_vars ≠ 0)
15      then x ← new NODE(_frames, NIL, INNER, _vars)
16      nodes.PUSH(x)
17    else x ← dd.GETNODE(GETSCENARIO(_frames))
18    x.frames ← x.frames ∪ _frames
19    n.ADDEDGE(v, x)
20  dd.INSERT(n)
21  n.ADDEDGE(OTHER, b)
22 dd.MERGESIMILARNODES()
23 for each n in dd.TRAVERSENODES()
24   do dd.TESTANDREMOVE(n)
25 for each n in dd.nodes
26   do n.REPLACEVALUEEDGESWITHINTERVALEDGE()
27 for each n in dd.nodes
28   do n.REORDEREDGES()
29 return dd
```

Fig. 14 The decision diagram construction algorithm.

nodes are assigned the empty set (line 3). A list with nodes that have to be processed is kept, and initially this list contains only the source node, unlabeled at this point (lines 5-6). While the list is not empty, the first node is extracted from it, and a variable that was not used on the path from the source to it is selected to label this node (lines 9-10). For each possible value for the selected variable that appears in the set of frames asso-

ciated with the node (line 12), an edge is added in the decision diagram (line 19). In line 13, the set of frames for which the prediction process goes through node n and for which the value of ξ matches v is saved. The new edge is added either to a new inner node that will go in the list of nodes to be processed (lines 15-16), or to a scenario node, in which case the list of frames of the scenario node is updated (lines 17-18). The decision is made in line 14 by checking if the list of variables that were not used for deciding the path from the source to the current node contains only the variable selected for labeling the currently processed node. Finally, the node is inserted into the decision diagram and an edge from it to the backup scenario node is created (lines 20-21).

2. *Node merging* (line 22): Two inner nodes are merged if they have the same label and the set of the outgoing edges of one is included in the set of the other one. To understand the reason behind this decision, consider the decision diagram of figure 13(a). It can be assumed that if $\xi_1 = 1$ and $\xi_3 = 4$ the application is, most probably, in scenario 2. This case did not appear for the training bitstream, but except for it the two ξ_3 labeled nodes imply the same decisions. If this assumption is made, the decision diagram can be reduced to the one shown in figure 13(b).
3. *Node removal* (lines 23-24): The diagram is traversed and each node is checked to see if it really influences the decision made by the diagram. If it does not, it can be removed. An example of this kind of node can be found in figure 13(b). In this diagram, it can be observed that whatever the values of ξ_1 and ξ_2 are, the current scenario is decided based on the value of ξ_3 (except for the values of ξ_1 and ξ_2 that did not occur in the training bitstream). This means that we can remove the nodes labeled with ξ_1 and ξ_2 from the diagram (see figure 13(c)). Note that if the values of ξ_1 and ξ_2 for a frame did not appear in the training bitstream, a scenario is selected based on the reduced diagram instead of the conservative backup scenario that would have been selected based on the original diagram.
4. *Interval edges* (lines 25-26): If a node has two or more outgoing edges associated to values $v_1 < v_2 < \dots < v_n$ that have the same destination, and there is no other outgoing edge associated with v , $v_1 < v < v_n$, then these edges may be merged in only one edge. In figure 13(c), for both $\xi_3 = 2$ and $\xi_3 = 4$, scenario 2 is selected and there is no other value for $\xi_3 \in [2, 4]$ for which another scenario is selected. The assumption that if a value $\xi_3 \in [2, 4]$ appears for a frame, scenario 2 should be selected with high probability, leads to the diagram figure 13(d).
5. *Edge reordering* (lines 27-28): To decrease the average runtime evaluation cost, the outgoing edges of each inner node are sorted in descending order based on the occur-

rence ratio of the values that label them. In figure 13(e), the edges for the node labeled with ξ_3 were reordered, based on the observation that $\xi_3 \in [2, 4]$ appears most often⁶.

Different optimization steps of our tool may be disabled, so the tool may produce different decision diagrams, from the one created only based on the training bitstream (only steps (1) and (5) of the above algorithm) to the one on which all possible size reductions were applied (all five steps). Also, in each step of the algorithm, for example, the selection of variables for labeling nodes (line 9), different heuristics may be used. However, it might be possible that by applying all steps the prediction quality becomes bad. This may happen as the decisions made in our diagram generation algorithm are based on practical observations, and the application at hand might not conform to these observations. In this case, the steps that negatively affect the prediction quality should be identified and disabled.

For each predictor, the average amount of cycles needed at runtime to predict the scenarios is profiled on the training bitstream and the scenario bounds are updated to accommodate for this prediction cost. The process is similar to the one used in the previous section for accommodating for the scenario switching cost.

In the experiments presented in section 8, we generated four fully optimized predictors, differentiated by:

- the variable selection heuristic for each node in step 1 of the algorithm (GETVAR, line 9 in figure 14): the variables with the most/least number of possible values are selected first. By selecting the one with most values first a lower runtime decision overhead might be introduced, as multiple small subtrees are created for each node and the decision height is reduced. On the other hand, by selecting the variable with the least possible values first, more freedom is given to the *interval edges* optimization step.
- the tree traversal in step 3 (TRAVERSENODE, line 23 in figure 14): breadth-/depth-first. Breadth-first tries to remove first the node, and then its children. Depth-first is doing the opposite.

All these four predictors can be used to achieve energy reduction, but there is no best one for all applications. Hence, in order to select the most efficient heuristics for an application, we generate the application source code for each of them. The structure of the generated source code is similar to the one presented in figure 15. It is derived from the original application, by introducing in it the predictor and the runtime quality preservation mechanism, which is described in the next section. Also, it contains the source code for adapting the processor frequency, which is activated only when

⁶ Scenario 2 from the decision diagram is the same as the scenario j computed in figure 9.

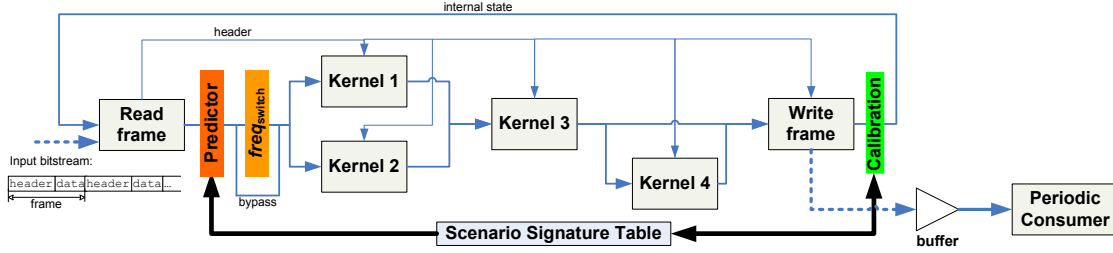


Fig. 15 Final implementation of the application.

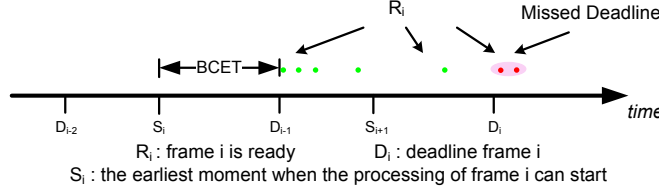


Fig. 16 Output buffer impact on processing start time.

the application switches from one scenario to another one. All the generated source codes are evaluated on the training bitstream and the one that gives the largest energy reduction is chosen. The variables used by its predictor are considered to be the most important control variables (fig. 6).

7 Quality preservation mechanism

Because of the variation in the time spent in processing a frame, usually, in real-time embedded systems, an output buffer is implemented (see the right part of figure 15). The smallest possible buffer has a size equal to the maximum size of a produced output frame. The buffer is used to avoid the stalling of the process until the periodic consumer (e.g. a screen) takes the produced frame, allowing the start of the next frame processing before the current frame is consumed. To implement this parallelism, the conflict situation of producing a new frame before the previous one was consumed should be handled. This can be done (i) by using a semaphore mechanism that postpones the writing until the frame is consumed, or (ii) by postponing the start moment of processing a new frame until it is sure that when the processing would be ready, the previous frame is already consumed.

We considered the second implementation, as there is no need for any synchronization mechanism. This gives more freedom in the consumer implementation and simplicity in output buffer implementation, for which a simple external memory may be used. Figure 16 explains how the start moment for frame processing is computed. For each frame i , S_i is defined as the earliest moment in time when the processing of frame i can start. It is equal to the moment when frame $i-1$ is consumed (D_{i-1}) minus the minimum possible processing time for each frame, estimated using static analysis

```

CALIBRATESCENARIOTABLE(INT scen,INT cycles)
1  framesCounter ++
2  if cycles > upperBound[scen]
3  then appMissesCounter ++
4  missCounter[scen] ++
5  maxBudget[scen] ← max(maxBudget[scen], cycles)
6  if appMissCounter / framesCounter > MISS-THRESHOLD
7  then s ← scen
8  for i ← 1 to noScenarios
9  do if missCounter[s] < missCounter[i]
10 then s ← i
11 UPDATESCENARIOINTERVAL(s, maxBudget[s])

```

Fig. 17 Runtime scenario quality control mechanism.

as the best case execution time (BCET) or measured. The proactive DVS-aware scheduler that we used in our experiments makes sure that a frame i does not start earlier than S_i . The processing of frame i can however also not start until frame $i-1$ is ready (R_{i-1}). If the deadline of frame $i-1$ is missed, so $R_{i-1} > D_{i-1}$, depending on the application, one of the following two decisions can be made: (i) the processing of frame $i-1$ might be stopped at D_{i-1} , so the processing of frame i can start or (ii) the application continues with the frame $i-1$ until it is ready, and then it starts with frame i . In the first case, which can for example be applied in an audio decoder, the processing of frame i actually starts at $\min(\max(S_i, R_{i-1}), D_{i-1})$. In the second one, typically used in video decoders that need a frame as a reference for the future, the processing of frame i starts at $\max(S_i, R_{i-1})$. For both ways to handle deadline misses, the consumer should not delete the frame from the output buffer when reading it, so it can read it again in case of a missed deadline. In our experiments from section 8 we consider the first case, as it fits the best with the selected benchmarks.

As in our approach the cycle budget required by the application for a specific frame is predicted based on the information collected on a training bitstream, it is possible

that the quality of the resulting system is lower than the required one, even when the above presented output buffer is exploited. This effect could appear because the training bitstream did not cover all the possible frames, so the scenario upper bounds might not be conservative. To keep the system quality under control, we introduce in the generated application source code a calibration mechanism (figure 15). This mechanism should be cheap in number of computation cycles and stored information size. We implemented it as in figure 17, and it adapts the table containing the scenario signatures used by the predictor. It counts the number of processed frames and the misses that appear in the system and for each scenario separately (lines 1-4). Also, for each scenario it stores (line 5) the maximum number of cycles that were used for processing a frame predicted to be in it. If the percentage of missed deadlines of the system is larger than a given threshold, the scenario with the largest number of misses is determined, and its cycle budget upper bound is updated (lines 6-11). As it was done also for the scenario switching mechanism and the predictor, the scenario bounds are updated to accommodate for the calibration mechanism too. Moreover, the overhead introduced by these three entities is taken into account when the cycle budget upper bound is updated at runtime (line 11).

8 Experimental results

All the steps of the presented tool-flow were implemented on top of SUIF [1], and they are applicable to applications written in C, as C is the most used language to write embedded systems software. The resulting implementation for the application is written in C, and has a structure similar to the one presented in figure 15.

We tested our method on three multimedia applications, an MP3 decoder [18], the motion compensation task of an MPEG-2 decoder [21] and a G.72x voice decompression algorithm [30]. The energy consumption was measured on an Intel XScale PXA255 processor [16], using the XTREM simulator [7]. We consider that the processor frequency (f_{CLK}) can be set discretely within the operational range of the processor, with 1MHz steps. The supply voltage (V_{DD}) is adapted accordingly, using the following equation:

$$f_{CLK} = k \cdot \frac{(V_{DD} - V_T)^2}{V_{DD}}, \quad (12)$$

where $V_T = 0.3V$ and the value of the constant k is computed for $V_{DD} = 1.5V$ and $f_{CLK} = 200MHz$. A frequency/voltage transition overhead $t_{switch} = 70\mu sec$ was considered, during which the processor stops running. The energy consumed during this transition is equal with $4\mu J$ [3]. When the processor is not used, it switches to an idle state within one

cycle, and it consumes an idle power of 63mW. This situation occurs if the start of a frame needs to be delayed, as explained in the previous section.

In the remaining part of this section, besides the main experiments that measure how much energy was saved by applying our approach, we quantify also the effect on energy of different steps of the decision diagram construction algorithm. Moreover, we investigate how the runtime calibration mechanism, different buffer sizes and different frequency/voltage switching costs influence the energy consumption and deadline miss ratio.

8.1 MP3 Decoder

The MPEG-I Layer III (MP3) decoder is a frame-based algorithm, which transforms a compressed bitstream in normal pulse code modulation data. A frame consists of 1152 mono or stereo frequency-domain samples, divided into two granules. The standard specifies a fixed decoding throughput: a frame at each 26ms. Details about the application structure and the source code are presented in [18]. To profile the application, we have chosen, as the training bitstream, a set of audio files consisting of: (i) the ones taken from [8], which were designed to cover all the extreme cases, and (ii) a few randomly selected stereo and mono songs downloaded from the internet, in order to cover the most common cases. After removing the data variables and loop iterators, the number of remaining control variables ξ_k to be considered for scenario prediction is 41. This set of variables is far more complete than the one detected using the static analysis from [13]. The scenario sets generation algorithm of section 5.3 leads to 2111 potential solutions (sets of scenarios). Using the method presented in section 5.4, we reduced the size of the pool of solutions for which the predictor was generated to 34. This decreases the execution time of the scenario analysis (section 6) from approximately 4 days to less than 5 hours. For each of the evaluated scenario sets, four fully optimized predictors were generated, as outlined in section 6.

To quantify the energy saved by our approach, we measured the energy consumed by the resulting application via three experiments, by decoding (i) 20 randomly selected stereo songs, (ii) 20 mono songs and (iii) all these 40 songs together. These three categories are the most common combinations of songs that appear during an MP3 decoder usage.

The three groups of bars of figure 18 present the normalized results of our approach, evaluated for two miss ratio thresholds as used in the calibration mechanism: 1% and 0.1%. The energy improvement is given relatively to the energy measured for the case when no scenarios knowledge was used. In this case, the frame cycle budget is the maximum number of cycles measured for all input frames. In each decoding period, first the frame is processed, and then

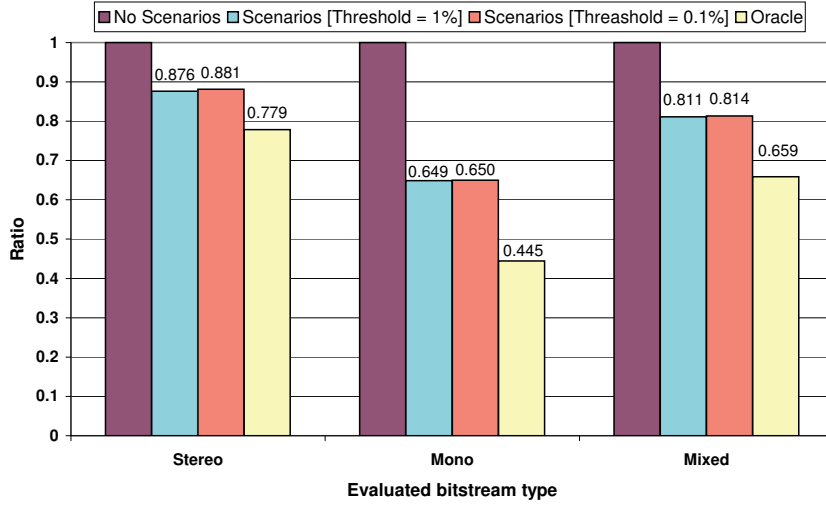


Fig. 18 Normalized energy consumption for the MP3 decoder.

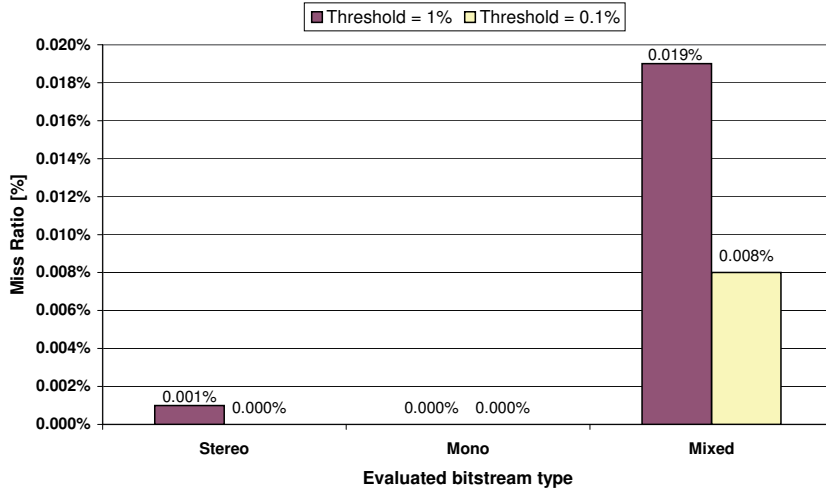


Fig. 19 Miss ratio for the MP3 decoder.

the processor goes in the idle state for the remaining time until the earliest possible start time for the next frame is reached.

We also compared our energy saving with the one given by an oracle (last bar of each group in figure 18), which is the smallest energy consumption that may be obtained. To compute it for a stream, all possible combinations of processor frequencies for decoding each frame from the stream were considered. The large difference between the energy reduction obtained by our approach and the oracle case is mostly due to the fact the oracle has a perfect knowledge of the remaining stream, based on which it may select different processor frequencies for the same scenario. Moreover, the oracle obtains an infinite accuracy without any cost, as it essentially considers any number of scenarios and variables for prediction, but has no prediction and calibration overhead. However, part of the energy difference is also due to the profiling drawbacks (e.g., not all possible samples were

covered) and due to the lack of a better scenario bound adaptation mechanism (e.g., a mechanism that allows the reduction of a scenario cycle upper bound). These problems may be overcome by using a more efficient runtime calibration algorithm that may also decrease the scenarios bounds and even modify the decision diagram. This topic is left for future work.

An important evaluation criterion for our approach is the percentage of missed deadlines. As the energy savings may lead to a miss ratio that is too high, we use a runtime calibration mechanism that allows us to set a threshold for the miss ratio. To evaluate the effectiveness of the calibration mechanism and the overall approach, we measured the miss ratio in the experiments. Figure 19 shows the results for the two selected thresholds. There is a relatively large difference between the imposed threshold and the measured miss ratio. This is because the threshold is constrained before the output buffer, and the miss ratio is measured after it. The out-

put buffer effect on miss ratio is hard to predict, but it will generally reduce the miss ratio. It can be observed that the combination of calibration and buffering is very effective.

Summarizing the main conclusions, for an MP3 player that is mainly used to listen mixed or stereo songs, the energy reduction that can be obtained by applying our approach is between 12% and 19%, for a miss ratio of up to one frame per 6 minutes (0.008%). The most energy efficient solution has 17 scenarios when decoding mixed (or only mono streams), and six when decoding only stereo streams.

Having concluded that our approach is effective, it is interesting to consider some of the design decisions in our approach, and some of the individual components in a bit more detail.

Recall that the decision diagram construction algorithm of section 6 uses two heuristics, one for labeling nodes in the diagram and one for traversing the diagram during the reduction. This leads to four possible combinations. For all three experiments we did, the most efficient predictor was the one generated by selecting during the decision diagram construction first the variables with the least number of possible values and by using a breadth-first reduction approach. This combination is the most effective one in many cases, although in some of our later experiments also other combinations turn out to be the most effective one.

To show that the runtime calibration mechanism and all the steps that we used during the decision diagram construction are relevant for energy reduction, we did eight different experiments for a threshold of 0.1% using the set of mixed streams as the benchmark, as shown in table 1. These experiments cover all possible cases for enabling/disabling three different components: (i) the runtime calibration mechanism, (ii) the node merging and removal (steps 2&3) in the decision diagram construction algorithm, and (iii) the usage of interval edges in the algorithm (step 4). The node merging and removal were considered together because they are very tightly linked: by merging some nodes, other nodes become irrelevant as decision makers, so they can be removed.

The most important observation from table 1 is that the merging and removal steps are essential to, and effective in, obtaining a substantial energy reduction. It turns out that when these optimization steps are omitted, 97% of the frames in the benchmark test falls into the backup scenario. This explains the low energy savings when the merging and removal steps are disabled. This also shows that the runtime prediction is not very effective in that case, which is in fact an indication that the training bitstream was not sufficiently representative to obtain a good predictor (without these optimizations). An important conclusion from these experiments is that the optimization steps in the decision diagram construction algorithm provide a high degree of robustness to our approach. They effectively resolved the shortcomings of a poor training bitstream. The results furthermore show

that also the interval optimization and the runtime calibration mechanism lead to further reductions in energy consumption. A final observation is that, for all the experiments, including the ones with the runtime calibration mechanism disabled, a set of scenarios and a predictor that meet the 0.1% miss ratio threshold was found. However, even if for this benchmark the required threshold could be met when the runtime calibration mechanism is not used, this will not be the case for all benchmarks and for all thresholds.

8.2 MPEG-2 Motion Compensation

An MPEG-2 video sequence is composed of frames, where each frame consists of a number of macroblocks (MBs). Decoding an MPEG-2 video can therefore be considered as decoding a sequence of MBs. This involves executing the following tasks for each MB: variable length decoding (VLD), inverse discrete cosine transformation (IDCT) and motion compensation (MC). Other tasks, like inverse quantization (IQ), involve a negligible amount of computation time, so we ignore them for the purpose of our analysis.

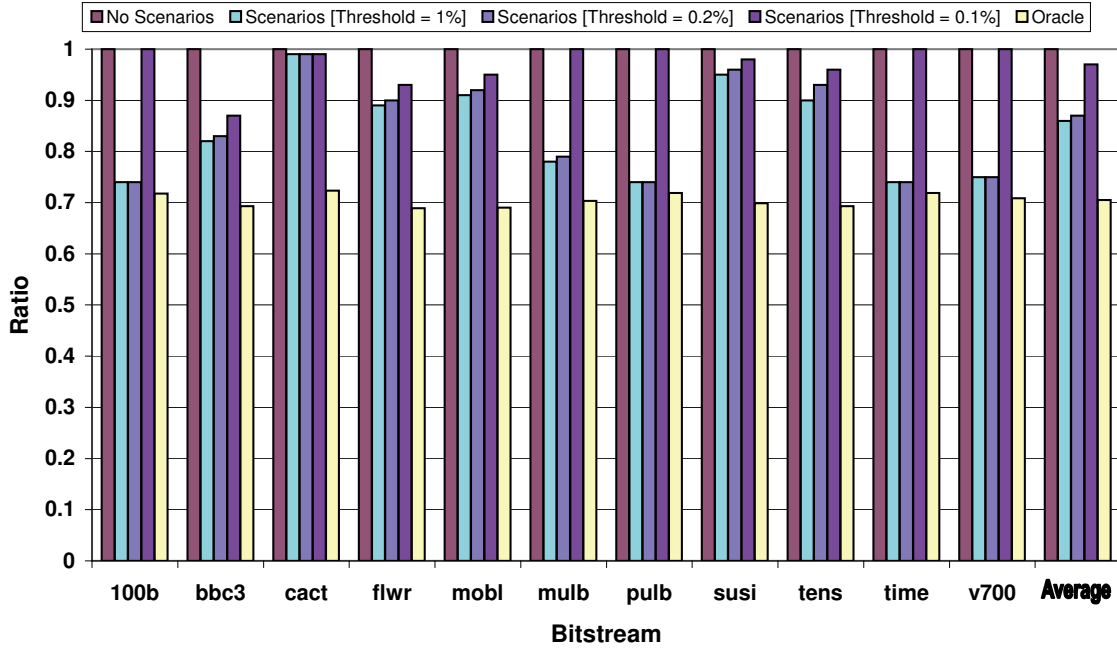
For our analysis, we use the source code from [21], and as a training bitstream we consider the first 20000 MBs from each test file from [31]. As the IDCT execution time for each MB is almost constant, we focus on MC and VLD. In case of the VLD, our tool could not discover the parameters that influence the execution time, as they do not exist in the code. This task is really data dependent, reading and processing the input stream for each MB until a stop flag is met. For the MC task, the parameters found by our tool include all the parameters identified manually in [2], and which can be found in the source code. Observe that when knowledge characterizing frame execution times is introduced in frame headers, as for example, proposed in [26], our tool will be able to fully automatically detect the variables that store this information, and then exploit it to obtain energy reductions.

In the remainder of the experiment, we focus on the MC task, for which the processing period of a MB is $120\mu sec$, which is very close to the frequency switching time $t_{switch} = 70\mu sec$. Therefore, we analyzed the possibility of using different values for the weight coefficient α in the cost function of equation (8). A larger value will give higher importance to reducing the number of runtime switches, than to reducing the over-estimation, and it will usually result in smaller scenario sets. We evaluated all α values between one and six, and we found that the best energy saving may be obtained for $\alpha = 3$.

The evaluation of our approach in terms of energy on the full streams of [31] is shown in figure 20. Three miss ratio thresholds were evaluated, the two used for the previous experiment (1% and 0.1%), and an intermediate one (0.2%). For this application, the most energy efficient solutions use three scenarios for the 1% and 0.2% miss ratio threshold,

Table 1 Experimental results for MP3 with a threshold of 0.1% miss ratio.

Decision diagram construction		Runtime calibration	#Scenarios	Selected predictor		Measured miss ratio	Energy reduction
Merging & Removal	Intervals			Var. selection	Reduction		
X	X	X	17	least values	breadth-first	0.008%	18.65%
X	-	X	17	least values	breadth-first	0.008%	15.46%
-	X	X	67	least values	-	0%	1.08%
-	-	X	67	least values	-	0%	1.08%
X	X	-	17	most values	breadth-first	0.085%	16.73%
X	-	-	17	least values	breadth-first	0.008%	15.46%
-	X	-	67	least values	-	0%	1.11%
-	-	-	67	least values	-	0%	1.08%

**Fig. 20** Normalized energy consumption for MPEG-2 MC.**Table 2** Experimental results for MPEG-2 MC with a threshold of 0.1% miss ratio.

Buffer size [macroblocks]	t_{switch} [μsec]	Energy reduction	Measured miss ratio
1	70	2.7%	0.029%
1	10	17.1%	0%
10	70	15.9%	0.008%

and two scenarios for the 0.1% threshold. The predictors were built by selecting, as for the MP3 decoder, first the variables with the least number of possible values, but using a depth-first instead of breadth-first reduction approach.

The measured miss ratio for all three thresholds is shown in figure 21. For a threshold of 0.2%, we obtained a 13% average energy reduction for all streams. The measured miss ratio was 0.09%, which represents one macroblock missed in every 13 frames when the video stream is in a QCIF format, that has a resolution of 176x144 pixels.

If the threshold is pushed to 0.1%, the energy reduction drops to 3%, as for three of the 11 streams, it was very difficult to obtain this miss ratio. This is due to the considered buffer that can accommodate only a variation in execution of at most $18\mu sec$, which is approximatively four times smaller than t_{switch} .

The results motivated us to do some experiments with varying buffer sizes and switching costs, to investigate their impact on energy savings and miss ratio. Table 2 shows the result of three experiments, the first one being the same experiment as reported in figures 20 and 21. It can be observed that a larger energy reduction for a 0.1% threshold (or any of the thresholds reported in figures 20 and 21) with a small measured miss ratio can be obtained when the frequency switching time t_{switch} is smaller or by increasing the output buffer size. The first might be obtained by using a different switching mechanism within the processor or another processor, and the second one is a viable solution when MC is considered in the context of a full MPEG-2 decoder. Then,

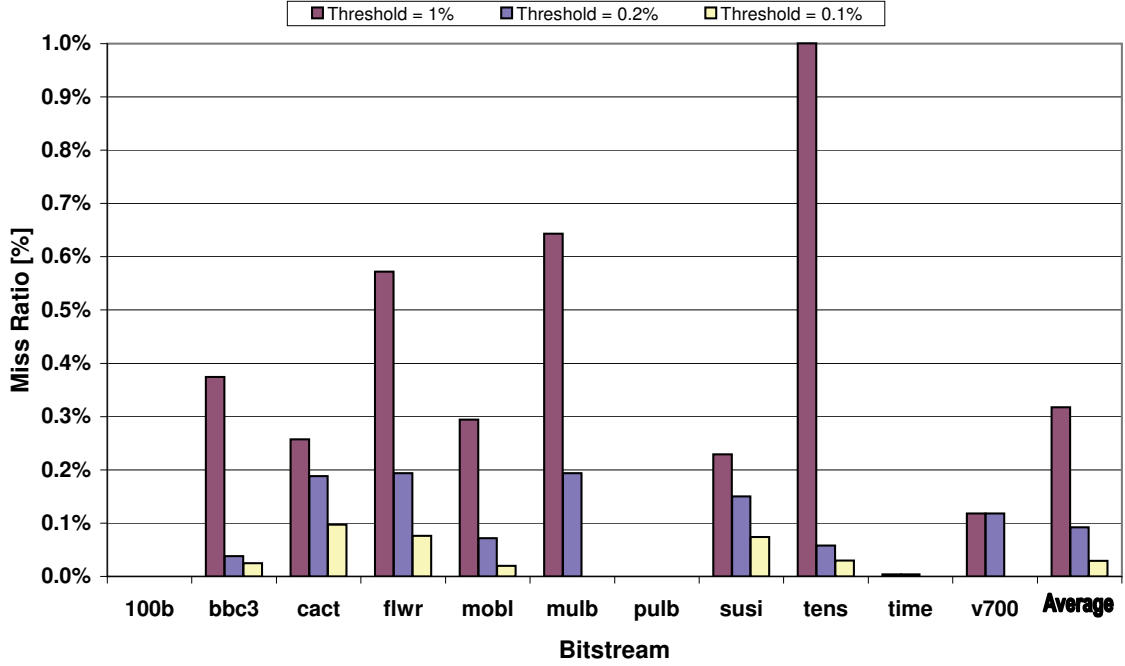


Fig. 21 Miss ratio for the MPEG-2 MC.

the buffer size can be increased without a supplementary cost, as the decoder already has to store the entire frame.

As a final remark, it should be noted that, when MC is embedded in a complete MPEG-2 decoder, the relative energy reduction observed by our approach will decrease. Even though MC is the most energy hungry component in the decoder, it does not count for more than 50% of the total energy. However, as already mentioned, if knowledge about frame execution times is introduced in the headers, as in [2, 15, 26], our tool will be able to exploit this information to optimize more components of the decoder.

8.3 G.72x Voice Decompression

This benchmark [30] implements the decoders for a set of G.721/G.723 adaptive differential pulse-code modulation (ADPCM) telephony speech codec standards covering the transmission of voice at rates of 24, 32, and 40 kbit/s. Its input streams are sampled at the rate of 8000 samples/second, so the deadline for each sample is $125\mu\text{sec}$.

We analyzed our approach on the streams of [6], using as training bitstream 3000 samples from each test file. The best energy saving was obtained using a set of three scenarios, each of them associated with a specific voice transmission rate: 24, 32 and 40 kbits/s. Figure 22 shows the results, both detailed per input type, and averaged. As for each stream the transmission rate is fixed, the number of runtime switches is exactly one, namely the initial scenario selection for the first sample from the stream. This, together with the fact that only one parameter is used in scenario detection, which helped in

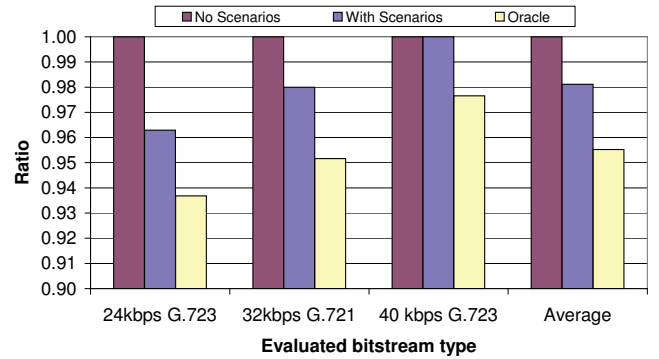


Fig. 22 Normalized energy consumption for the G.72X voice decompression.

having a fully representative training bitstream, leads to a miss ratio equal to zero for any imposed threshold. So, even if the resulting improvement is small (just 2%), it comes for free, without quality reduction. Furthermore, our method realizes close to 50% of the maximum theoretical possible improvement of slightly over 4%, computed via the oracle. The result of almost 50% of the theoretical maximum is inline with the earlier two experiments.

9 Conclusion

In this paper, we have presented a profiling driven approach to detect and characterize scenarios for single-task soft real-time multimedia applications. The scenarios are identified based on the automatically detected control variables whose values influence the application execution time the most. In

addition, we present a technique to automatically derive and insert predictors in the application code, which are used at runtime to select the current scenario. Our method is fully automated and it was tested on three multimedia applications. For all of them, the identified sets of variables are similar to manually selected sets. We show that, using a proactive DVS-aware scheduler based on the scenarios and the runtime predictor generated by our tool using the identified variables, energy consumption decreases with up to 19%, having guaranteed, using a simple runtime calibration mechanism, a frame deadline miss ratio of less than 0.1%. In practice, due to output buffering, the measured miss ratio decreases even to almost zero.

In future work, we would like to investigate different runtime calibration algorithms, that learn on the fly and adapt the scenario bounds, the number of scenarios and the decision diagram underlying the predictor. The information collected and processed by these control algorithms will be used not only for keeping the miss ratio under control, but also for further reduction in energy consumption. We also plan to extend our work to multi-task applications. Even if most of the basic steps of the presented trajectory (e.g., parameter identification, scenario prediction) remain unchanged, others, particularly scenario selection, have to be adapted to accommodate the specific problems that appear in multi-task applications (e.g., communication delay between tasks, pipelined execution). Moreover, scenario based design is not limited to multimedia applications and execution time estimation. It is interesting to investigate to what extent our techniques can be applied to other systems and/or other resource costs (such as memory accesses). Again, parameter identification and scenario prediction seem relatively straightforward to adapt. Scenario selection is the step that depends the most on the particular context.

References

1. Amarasinghe, S.P., Anderson, J.M., Lam, M.S., Lim, A.W.: An overview of a compiler for scalable parallel machines. In: Proc. of the 6th International Workshop on Languages and Compilers for Parallel Computing, pp. 253–272. Springer-Verlag, Germany (1993)
2. Bavier, A.C., Montz, A.B., Peterson, L.L.: Predicting MPEG execution times. *ACM SIGMETRICS Performance Evaluation Review* **26**(1), 131–140 (1998)
3. Burd, T.D., Pering, T.A., Stratakos, A.J., Brodersen, R.W.: A dynamic voltage scaled microprocessor system. *IEEE Journal of Solid-State Circuits* **35**(11), 1571–1580 (2000)
4. Carroll, J.M. (ed.): Scenario-based design: envisioning work and technology in system development. John Wiley & Sons Inc, New York, NY (1995)
5. Chawathe, S.S., Rajaraman, A., Garcia-Molina, H., Widom, J.: Change detection in hierarchically structured information. *ACM SIGMOD Record* **25**(2), 493–504 (1996)
6. Clamen, S.M.: 8bit ULAW files collection (2006). <http://www.cs.cmu.edu/People/clamen/misc/tv/Animaniacs/sounds/>
7. Contreras, G., Martonosi, M., Peng, J., Ju, R., Lueh, G.Y.: XTREM: A power simulator for the Intel XScale core. *ACM SIGPLAN Notices* **39**(7), 115–125 (2004)
8. Dietz, M., et al.: MPEG-1 audio layer III test bitstream package (1994). <http://www.iis.fhg.de>
9. Douglass, B.: Real Time UML: Advances in the UML for Real-Time Systems. Addison Wesley Publishing Company, Reading, MA (2004)
10. Gheorghita, S.V., Basten, T., Corporaal, H.: Intra-task scenario-aware voltage scheduling. In: Proc. of the International Conference on Compilers, Architecture and Synthesis for Embedded Systems (CASES), pp. 177–184. ACM Press, New York, NY (2005)
11. Gheorghita, S.V., Basten, T., Corporaal, H.: Application scenarios in streaming-oriented embedded system design. In: Proc. of the International Symposium on System-on-Chip (SoC 2006), pp. 175–178. IEEE Computer Society Press, Los Alamitos, CA (2006)
12. Gheorghita, S.V., Basten, T., Corporaal, H.: Profiling driven scenario detection and prediction for multimedia applications. In: Proc. of the International Conference on Embedded Computer Systems: Architectures, Modeling, and Simulation (IC-SAMOS), pp. 63–70. IEEE Computer Society Press, Los Alamitos, CA (2006)
13. Gheorghita, S.V., Stuijk, S., Basten, T., Corporaal, H.: Automatic scenario detection for improved WCET estimation. In: Proc. of the 42nd Design Automation Conference DAC, pp. 101–104. ACM Press, New York, NY (2005)
14. Hind, M., Burke, M., Carini, P., Choi, J.: Interprocedural pointer alias analysis. *ACM Transactions on Programming Languages and Systems* **21**(4), 848–894 (1999)
15. Huang, Y., Chakraborty, S., Wang, Y.: Using offline bitstream analysis for power-aware video decoding in portable devices. In: Proc. of the 13th ACM International Conference on Multimedia, pp. 299–302. ACM Press, New York, NY (2005)
16. Intel Corporation: Intel XScale microarchitecture for the PXA255 processor: Users manual (2003). Order No. 278796
17. Jha, N.K.: Low power system scheduling and synthesis. In: Proc. of the IEEE/ACM International Conference on Computer Aided Design (ICCAD), pp. 259–263. IEEE Computer Society Press, Los Alamitos, CA (2001)
18. Lagerström, K.: Design and implementation of an MP3 decoder (2001). URL <http://www.kmlager.com/mp3/>. M.Sc. thesis, Chalmers University of Technology, Sweden
19. Maxiaguine, A., Liu, Y., Chakraborty, S., Ooi, W.T.: Identifying “representative” workloads in designing MpSoC platforms for media processing. In: Proc. of 2nd Workshop on Embedded Systems for Real-Time Multimedia (ESTIMedia), pp. 41–46. IEEE Computer Society Press, Los Alamitos, CA (2004)
20. McCluskey, E.J.: Minimization of boolean functions. *Bell System Technical Journal* **35**(5), 1417–1444 (1956)
21. MPEG Software Simulation Group: MPEG-2 video codec (2006). ftp://ftp.mpegv.com/pub/mpeg/mssg/mpeg2vidcodec_v12.tar.gz
22. Murali, S., Coenen, M., Radulescu, A., Goossens, K., DeMicheli, G.: A methodology for mapping multiple use-cases onto networks on chips. In: Proc. of Design, Automation, and Test in Europe (DATE), pp. 118–123. IEEE Computer Society Press, Los Alamitos, CA (2006)
23. Palkovic, M., Corporaal, H., Catthoor, F.: Global memory optimisation for embedded systems allowed by code duplication. In: Proc. of the 9th International Workshop on Software and Compilers for Embedded Systems (SCOPEs), pp. 72–79. ACM Press, New York, NY (2005)
24. Paul, J.M., Thomas, D.E., Bobrek, A.: Scenario-oriented design for single-chip heterogeneous multiprocessors. *IEEE Transactions on Very Large Scale Integration (VLSI) Systems* **14**(8), 868–880 (2006)

25. Pedram, M., Cheng, W.C., Dantu, K., Choi, K.: Frame-based dynamic voltage and frequency scaling for a MPEG decoder. In: Proc. of IEEE/ACM International Conference on Computer-Aided Design (ICCAD), pp. 732–737. ACM Press, New York, NY (2002)
26. Poplavko, P., Basten, T., Pastnak, M., van Meerbergen, J., Bekooij, M., de With, P.: Estimation of execution times of on-chip multiprocessors stream-oriented applications. In: Proc. of the 3rd ACM/IEEE International Conference in Formal Methods and Models for Codesign (MEMOCODE), pp. 251–252. IEEE Computer Society Press, Los Alamitos, CA (2005)
27. Rutten, M.J., van Eijndhoven, J.T.J., Jaspers, E.G.T., van der Wolf, P., Pol, E.D., Gangwal, O.P., Timmer, A.: A heterogeneous multiprocessor architecture for flexible media processing. *IEEE Design & Test of Computers* **19**(4), 39–50 (2002)
28. Sasanka, R., Hughes, C.J., Adve, S.V.: Joint local and global hardware adaptations for energy. *ACM SIGARCH Computer Architecture News* **30**(5), 144–155 (2002)
29. Shin, D., Kim, J.: Optimizing intra-task voltage scheduling using data flow analysis. In: Proc. of the 10th Asia and South Pacific Design Automation Conference (ASP-DAC), pp. 703–708. ACM Press, New York, NY (2005)
30. Sun Microsystems, Inc.: Free implementation of CCITT compression types G.711, G.721 and G.723 (2006)
31. Tektronix: MPEG-2 video test bitstreams (2006). <ftp://ftp.tek.com/tv/test/streams/Element/MPEG-Video/525/>
32. Wegener, I.: Integer-Valued DDs. In: *Branching Programs and Binary Decision Diagrams: Theory and Applications*, SIAM Monographs on Discrete Mathematics and Applications, chap. 9. Society for Industrial and Applied Mathematics, Philadelphia, PA (2000)
33. Yang, P., Marchal, P., Wong, C., Himpe, S., Catthoor, F., David, P., Vounckx, J., Lauwereins, R.: Cost-efficient mapping of dynamic concurrent tasks in embedded real-time multimedia systems. In: W. Wolf, A. Jerraya (eds.) *Multi-Processor Systems on Chip*, chap. 11. Morgan Kaufmann Publishers, San Francisco, CA (2003)