

Hartmann, Sönke

Working Paper — Digitized Version

A comparative genetic algorithm for resource-constraint project scheduling

Manuskripte aus den Instituten für Betriebswirtschaftslehre der Universität Kiel, No. 451

Provided in Cooperation with:

Christian-Albrechts-University of Kiel, Institute of Business Administration

Suggested Citation: Hartmann, Sönke (1997) : A comparative genetic algorithm for resource-constraint project scheduling, Manuskripte aus den Instituten für Betriebswirtschaftslehre der Universität Kiel, No. 451, Universität Kiel, Institut für Betriebswirtschaftslehre, Kiel

This Version is available at:

<https://hdl.handle.net/10419/177312>

Standard-Nutzungsbedingungen:

Die Dokumente auf EconStor dürfen zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden.

Sie dürfen die Dokumente nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, öffentlich zugänglich machen, vertreiben oder anderweitig nutzen.

Sofern die Verfasser die Dokumente unter Open-Content-Lizenzen (insbesondere CC-Lizenzen) zur Verfügung gestellt haben sollten, gelten abweichend von diesen Nutzungsbedingungen die in der dort genannten Lizenz gewährten Nutzungsrechte.

Terms of use:

Documents in EconStor may be saved and copied for your personal and scholarly purposes.

You are not to copy documents for public or commercial purposes, to exhibit the documents publicly, to make them publicly available on the internet, or to distribute or otherwise use the documents in public.

If the documents have been made available under an Open Content Licence (especially Creative Commons Licences), you may exercise further usage rights as specified in the indicated licence.

Manuskripte
aus den
Instituten für Betriebswirtschaftslehre
der Universität Kiel

No. 451

**A Competitive Genetic Algorithm for
Resource-Constrained Project Scheduling**

Sönke Hartmann

August 1997

Sönke Hartmann¹
Institut für Betriebswirtschaftslehre
Lehrstuhl für Produktion und Logistik
Christian-Albrechts-Universität zu Kiel
Olshausenstraße 40
24118 Kiel, Germany

e-mail: hartmann@bwl.uni-kiel.de

FTP: <ftp://ftp.bwl.uni-kiel.de/pub/operations-research/>

WWW: <http://www.bwl.uni-kiel.de/bwlinstitute/Prod/>

¹supported by the *Studienstiftung des deutschen Volkes*

Abstract

In this paper we consider the resource-constrained project scheduling problem (RCPSP) with makespan minimization as objective. We propose a new genetic algorithm approach to solve this problem. Subsequently, we compare it to two genetic algorithm concepts from the literature. While our approach makes use of a permutation based genetic encoding that contains problem-specific knowledge, the other two procedures employ a priority value based and a priority rule based representation, respectively. Then we present the results of our thorough computational study for which a standard set of project instances has been used. The outcome reveals that our procedure is the most promising genetic algorithm to solve the RCPSP. Finally, a priority rule based random sampling procedure known from the literature serves as a further benchmark. We show that our genetic algorithm yields better results than this sampling approach.

Keywords: Project Management and Scheduling, Resource-Constraints, Genetic Algorithms, Computational Results.

1 Introduction

Within the classical resource-constrained project scheduling problem (RCPSP), the activities of a project have to be scheduled such that the makespan of the project is minimized. This problem arises within project management software as well as within systems for production planning and scheduling. The currently most powerful exact procedures have been presented by Brucker et al. [2], Demeulemeester and Herroelen [5, 6], Mingozzi et al. [22], and Sprecher [26]. However, they are unable to find optimal schedules for highly resource-constrained projects with 60 activities or more. Hence, in practice heuristic algorithms to generate near-optimal schedules for larger projects are of special interest. Recently, an evaluation study (see Kolisch and Hempel [16] and Kolisch [15]) showed that commercial software packages for project management generate schedules with an average deviation of 4.3 % to 9.8 % from the optimal solution for projects with up to 30 activities. These rather disappointing results indicate the need for (fast) heuristics to obtain better near-optimal schedules which should be implemented in software packages.

Recently proposed heuristic algorithms for the RCPSP include the following: Kolisch [13] develops three new priority rules for the parallel scheduling scheme and tests them in a single-pass environment. Kolisch [14] compares some good priority rules within both the serial and the parallel scheduling scheme and reports experiences with sampling procedures. Kolisch and Drexel [17] introduce a so-called adaptive search procedure which applies a scheduling scheme and priority rules after analyzing the project instance at hand. Baar et al. [1] develop two tabu search approaches. Lee and Kim [20] compare a genetic algorithm (GA), a simulated annealing heuristic, and a tabu search method. Cho and Kim [3] improve the simulated annealing procedure of Lee and Kim [20]. Kohlmorgen et al. [12] summarize their experiences with an implementation of a GA for the RCPSP on parallel processors. For the more general RCPSP with multiple execution modes (MRCPPSP), GAs have been proposed by Özdamar [24], Hartmann [8], and Mori and Tseng [23].

The purpose of this paper is to introduce a new GA approach for solving the RCPSP and to compare it to two existing GA concepts for this problem class. The three procedures are based on different genetic encodings and encoding-specific genetic operators. We proceed

as follows: Section 2 provides a description of the RCPSP. Section 3 introduces the new GA which makes use of a permutation based representation. Sections 4 and 5 summarize two GA approaches from the literature which employ a priority value and a priority rule based encoding, respectively. Section 6 gives the results of our in-depth computational experiments. Finally, Section 7 states some conclusions.

2 Problem Description

We consider a project which consists of J activities (jobs) labeled $j = 1, \dots, J$. Due to technological requirements, there are precedence relations between some of the jobs. These precedence relations are given by sets of immediate predecessors P_j indicating that an activity j may not be started before all of its predecessors are completed. Analogously, S_j is the set of the immediate successors of activity j . The transitive closure of the precedence relations is given by sets of (not necessarily immediate) successors $\bar{S}_j$. The precedence relations can be represented by an activity-on-node network which is assumed to be acyclic. We consider additional activities $j = 0$ representing the single source and $j = J + 1$ representing the single sink activity of the network.

With the exception of the (dummy) source and (dummy) sink activity, each activity requires certain amounts of (renewable) resources to be performed. The set of resources is referred to as K . For each resource $k \in K$ the per-period-availability is constant and given by R_k . The processing time (duration) of an activity j is denoted as p_j , its request for resource k is given by r_{jk} . Once started, an activity may not be interrupted. W.l.o.g., we assume that the dummy source and the dummy sink activity have a duration of zero periods and no request for any resource.

The parameters are assumed to be nonnegative and integer valued. The objective is to determine a schedule with minimal makespan such that both the precedence and resource constraints are fulfilled. Mathematical programming formulations of the RCPSP have been given by e.g. Demeulemeester and Herroelen [5, 6], Mingozi et al. [22], and Sprecher [26].

3 Permutation based Genetic Algorithm

3.1 Basic Scheme

Introduced by Holland [11], genetic algorithms (GAs) serve as a heuristic meta strategy to solve hard optimization problems. For an introduction into GAs, we refer to Goldberg [7]. In this section we propose a new GA approach for the RCPSP.

The GA starts by computing an initial population, i.e. the first generation. We assume that the initial population contains POP individuals where POP is an even integer. After computing the fitness values of the individuals, the population is randomly partitioned into pairs of individuals. To each resulting pair of (parent) individuals, we apply the crossover operator to produce two new (children) individuals. Subsequently, we apply the mutation operator to the genotypes of the newly produced children. After determining the fitness of each child individual, we add the children to the current population, leading to a population size of $2 \cdot POP$. Finally, applying the selection operator to reduce the population to its former size, we obtain the next generation. The algorithm stops if a prespecified number

of generations which is denoted as GEN or a given time limit is reached. Clearly, at most $POP \cdot GEN$ different individuals (and related schedules) are calculated.

3.2 Individuals and fitness

In the first GA variant to be examined, an individual I is represented by an activity sequence $j_1^I, \dots, j_J^I$. This job sequence is assumed to be a precedence feasible permutation of the set of activities, that is, we have $\{j_1^I, \dots, j_J^I\} = \{1, \dots, J\}$ and $P_{j_i^I} \subseteq \{0, j_1^I, \dots, j_{i-1}^I\}$ for $i = 1, \dots, J$.

Each genotype is related to a uniquely determined schedule (phenotype) which is computed using the following serial scheduling scheme: First, the dummy source activity is started at time 0. Then we schedule the activities in the order that is prescribed by the sequence $j_1^I, \dots, j_J^I$. Thereby, each activity is assigned the earliest feasible start time. Note that the result is an active schedule, cf. Kolisch [14]. That is, no activity can be left shifted without violating the constraints (for a formal definition of active schedules cf. Sprecher et al. [27]). The fitness of an individual I is given by the makespan of the related schedule. Consider the project instance shown in Figure 1 and the individual displayed in Figure 2 (a). Applying the serial scheduling procedure described above leads to the schedule given in Figure 3. The fitness of this individual is 13.

The initial population is computed as follows: Starting with the empty job sequence, we obtain a precedence feasible job sequence by repeatedly applying the following step: The next activity in the job sequence is randomly taken from the set of those currently unselected activities all non-dummy predecessors of which have already been selected for the job sequence. In addition to this random algorithm, we have tested another variant which determines the initial population with a sampling procedure as described by Kolisch [14]. More precisely, we employ a good priority rule, in our case the latest finish time rule (LFT), to derive probabilities which are used to select the next activity for the job sequence.

Notice that, while each individual is related to a unique schedule, a schedule can be related to more than one individual. In other words, there is some redundancy in the search space as distinct elements of the search space (i.e. genotypes) may be related to the same schedule. Consider again the project instance of Figure 1 and the genotype of Figure 2 (a). Obviously, we obtain a different genotype by exchanging activities 1 and 6 in the job sequence. However, both genotypes are related to the same schedule, i.e. the schedule displayed in Figure 3.

3.3 Crossover

We consider three different crossover variants for the permutation based encoding. They are similar to the general crossover technique described by Reeves [25] for permutation based genotypes, with the difference that our encoding takes precedence relations into account.

The first crossover operator is called **one-point crossover**. We consider two individuals selected for crossover, a mother M and a father F . Then we draw a random integer q with $1 \leq q < J$. Now two new individuals, a daughter D and a son S , are produced from the parents. We first consider D which is defined as follows: In the job sequence of D , the positions $i = 1, \dots, q$ are taken from the mother, that is,

$$j_i^D := j_i^M.$$

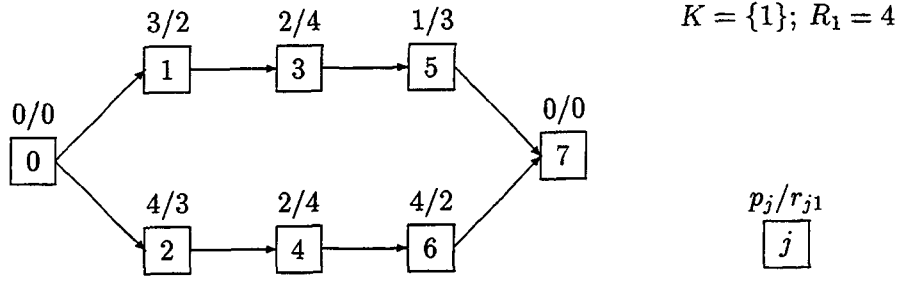


Figure 1: Project instance

(a) permutation	2	4	1	6	3	5
(b) priority value	0.58	0.64	0.31	0.87	0.09	0.34
(c) priority rule	LST	GRPW	MTS	LST	MSLK	LFT

Figure 2: Example individuals

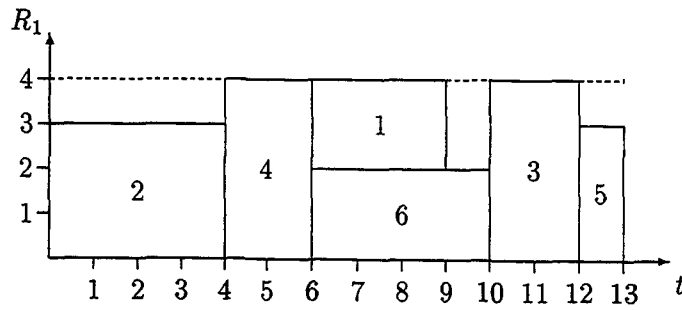


Figure 3: Example schedule

The job sequence of positions $i = q + 1, \dots, J$ in D is taken from the father. However, the jobs that have already been taken from the mother may not be considered again. We obtain

$$j_i^D := j_k^F \text{ where } k \text{ is the lowest index such that } j_k^F \notin \{j_1^D, \dots, j_{i-1}^D\}.$$

As a result, the relative positions in the parents' job sequences are preserved. While this definition obviously ensures that each activity appears exactly once in the resulting job sequence, the following theorem shows that also the precedence assumption is fulfilled.

Theorem 1 *If applied to precedence feasible parent individuals, the one-point crossover operator for the permutation based genetic encoding results in a precedence feasible offspring genotype.*

Proof. Let the genotypes of the parents M and F fulfill the precedence assumption. We assume that the child individual D produced by the crossover operator is not precedence feasible. That is, there are two activities j_i^D and j_k^D with $1 \leq i < k \leq J$ and $j_k^D \in P_{j_i^D}$. Three cases can be distinguished:

Case 1: We have $i, k \leq q$. Then activity j_i^D is before activity j_k^D in the job sequence of M , a contradiction to the precedence feasibility of M .

Case 2: We have $i, k > q$. As the relative positions are maintained by the crossover operator, activity j_i^D is before activity j_k^D in the job sequence of F , contradicting the precedence feasibility of F .

Case 3: We have $i \leq q$ and $k > q$. Then activity j_i^D is before activity j_k^D in the job sequence of M , again a contradiction to the precedence feasibility of M . $\square$

The son S of the individuals M and F is computed similarly. However, the positions $1, \dots, q$ of the son's job sequence are taken from the father and the remaining positions are determined by the mother. Obviously, Theorem 1 holds for both offspring individuals D and S .

The second crossover operator is an extension of the one-point variant, called **two-point crossover**. Here, we draw two random integers q_1 and q_2 with $1 \leq q_1 < q_2 \leq J$. Now the daughter individual D is determined by taking the job sequence of the positions $i = 1, \dots, q_1$ from the mother, that is,

$$j_i^D := j_i^M.$$

The positions $i = q_1 + 1, \dots, q_2$ are derived from the father:

$$j_i^D := j_k^F \text{ where } k \text{ is the lowest index such that } j_k^F \notin \{j_1^D, \dots, j_{i-1}^D\}.$$

The remaining positions $i = q_2 + 1, \dots, J$ are again taken from the mother, that is,

$$j_i^D := j_k^M \text{ where } k \text{ is the lowest index such that } j_k^M \notin \{j_1^D, \dots, j_{i-1}^D\}.$$

The son individual is computed analogously, taking the first and third part from the father and the second one from the mother. Obviously, Theorem 1 can easily be extended to the two-point crossover. Observe also that fixing $q_2 = J$ leads to the one-point variant which therefore is a special case of the two-point crossover.

The third crossover type is called **uniform crossover**. Here, the daughter D is determined as follows: We draw a sequence of random numbers $p_i \in \{0, 1\}$, $i = 1, \dots, J$. For each position $i = 1, \dots, J$ we take the activity from the mother, if we have $p_i = 1$, that is,

$$j_i^D := j_k^M \text{ where } k \text{ is the lowest index such that } j_k^M \notin \{j_1^D, \dots, j_{i-1}^D\}.$$

Otherwise, if $p_i = 0$, the activity is derived from the father's job sequence:

$$j_i^D := j_k^F \text{ where } k \text{ is the lowest index such that } j_k^F \notin \{j_1^D, \dots, j_{i-1}^D\}.$$

The son S is computed using an analogous procedure which takes the i -th job from the father if $p_i = 1$ and from the mother otherwise. Note that the uniform crossover generalizes the two-point variant: Fixing $p_i = 1$ for $i \in \{1, \dots, q_1, q_2 + 1, \dots, J\}$ and $p_i = 0$ for $i \in \{q_1 + 1, \dots, q_2\}$ leads to the definition of the daughter in the two-point crossover. With arguments similar to those used in the proof of Theorem 1, one can show that also the uniform crossover produces precedence feasible offspring.

3.4 Mutation

Given a permutation based individual I , the mutation operator modifies the related job sequence as follows: For all positions $i = 1, \dots, J - 1$, activities j_i^I and j_{i+1}^I are exchanged with a probability of p_{mutation} , if the result is a job sequence which fulfills the precedence assumption.

The mutation operator may create job sequences (i.e. gene combinations) that could not have been produced by the crossover operator. However, it should be noted that performing a mutation on an individual does not necessarily change the related schedule. This is due to the redundancy in the genetic representation: For example, interchanging two activities in the job sequence which have the same start time changes the individual, but not the related schedule.

3.5 Selection

We consider four alternative selection operators which follow a survival-of-the-fittest strategy as similarly described by e.g. Michalewicz [21]. The first one is a simple **ranking method**: We keep the POP best individuals and remove the remaining ones from the population (ties are broken arbitrarily).

The second variant, the **proportional selection**, can be viewed as a randomized version of the previously described ranking technique. Let $f(I)$ be the fitness of an individual I , and let $\mathcal{P}$ denote the current population, that is, a list containing the individuals. Note that we use a list of individuals instead of a set because we explicitly allow two (or more) distinct individuals with the same genotype in a population. We restore the original population size by successively removing individuals from the population until POP individuals are left, using the following probability: Denoting with $f_{\text{best}} = \min\{f(I) \mid I \in \mathcal{P}\}$ the best fitness in the current population, the probability to die for an individual I is given by

$$p_{\text{death}}(I) = \frac{(f(I) - f_{\text{best}} + 1)^2}{\sum_{I' \in \mathcal{P}} (f(I') - f_{\text{best}} + 1)^2}.$$

Next, we consider two versions of the tournament technique. In the **2-tournament selection**, two different randomly chosen individuals I_1 and I_2 compete for (temporary) survival. If individual I_1 is not better than individual I_2 , i.e. $f(I_1) \geq f(I_2)$, then it dies and is removed from the population (again, ties are broken arbitrarily). This process is repeated until *POP* individuals are left.

Finally, the **3-tournament selection** extends the previously described method by randomly selecting three individuals I_1 , I_2 , and I_3 . If we have $f(I_1) \geq f(I_2)$ and $f(I_1) \geq f(I_3)$, individual I_1 is removed from the population. Again, this step is repeated until *POP* individuals are left.

4 Priority Value based Genetic Algorithm

This section is devoted to a GA which makes use of a priority value based encoding as similarly used by Lee and Kim [20] within their GA for the RCPSP. We employ the priority value representation into the basic GA scheme that was used for the permutation based GA. This allows us to obtain comparable results when evaluating the GA approaches. We also make use of the selection operator variants here that have been defined for the permutation based GA as they are not encoding specific. In the following, the priority value based representation and the related crossover and mutation operators are discussed.

4.1 Individuals and fitness

In this GA approach, an individual I is represented by a sequence of priority values $I = pv_1^I, \dots, pv_J^I$. For each priority value of activity $j = 1, \dots, J$, we have $pv_j^I \in [0, 1]$. This encoding equals the one employed by Lee and Kim [20]. In contrast to their approach which employs a parallel scheduling scheme to derive the schedule related to an individual, however, we employ a serial scheme similar to that used for the permutation based encoding. We do so because it has been shown by Kolisch [14] that a serial scheme yields better results when a large number of schedules is computed for one project instance. This is due to the fact that using the parallel scheme, one might exclude all optimal solutions from the search space while the search space of the serial scheme always contains an optimal schedule.

Hence, given an individual I , the related schedule is computed as follows: After starting the dummy source activity at time 0, we determine the set EJ of the eligible activities, that is, those activities the predecessors of which are already scheduled. Then we schedule the eligible activity j with the highest priority value $pv_j^I = \max \{pv_i^I \mid i \in EJ\}$ as early as possible such that neither the precedence nor the resource constraints are violated. Repeatedly scheduling an unscheduled activity, we obtain a feasible active schedule. Again, the fitness of an individual is defined as the makespan of the related schedule. Consider again the instance of Figure 1. The example individual displayed in Figure 2 (b) is related to the schedule of Figure 3.

Each individual I of the initial population is determined by randomly drawing a priority value $pv_j^I \in [0, 1]$ with a uniform distribution for each activity $j = 1, \dots, J$.

As was the case for the permutation based encoding, there is some redundancy in the search space also for the priority value based representation. We consider again the individual shown in Figure 2 (b). Clearly, setting for example $pv_2^I = 0.93$ instead of 0.64, we

obtain a different individual. However, both individuals are related to the same schedule, namely the one of Figure 3.

4.2 Crossover

This encoding allows us to employ standard crossover operators. Again, we consider two individuals selected for crossover, a mother M and a father F , from which two offspring individuals are computed. In the following, we only define the daughter D . As for the permutation based representation, the son S is computed analogously to the daughter's definition.

For the **one-point crossover**, we draw a random integer q with $1 \leq q < J$. The first q positions of daughter individual D are taken from the mother while the remaining ones are defined by the father, that is, for each $i = 1, \dots, J$ we have

$$pv_i^D = \begin{cases} pv_i^M, & \text{if } i \in \{1, \dots, q\} \\ pv_i^F, & \text{if } i \in \{q+1, \dots, J\}. \end{cases}$$

Similarly to the permutation based encoding, we have considered the **two-point crossover** in addition to the one-point variant defined above. We draw two random integers q_1 and q_2 with $1 \leq q_1 < q_2 \leq J$ and obtain for each $i = 1, \dots, J$

$$pv_i^D = \begin{cases} pv_i^M, & \text{if } i \in \{1, \dots, q_1\} \\ pv_i^F, & \text{if } i \in \{q_1+1, \dots, q_2\} \\ pv_i^M, & \text{if } i \in \{q_2+1, \dots, J\}. \end{cases}$$

Finally, for the **uniform crossover**, we draw a sequence of random numbers $p_i \in \{0, 1\}$, $i = 1, \dots, J$. Then we set for each $i = 1, \dots, J$

$$pv_i^D = \begin{cases} pv_i^M, & \text{if } p_i = 1 \\ pv_i^F, & \text{otherwise.} \end{cases}$$

4.3 Mutation

The mutation for the priority value based encoding is defined as follows: Given an individual I , we modify the related priority value sequence as follows: For all positions $i = 1, \dots, J$, a new priority value $pv_i^I \in [0, 1]$ is drawn with a probability of p_{mutation} . Clearly, the result is always a feasible priority value sequence. Obviously, the mutation operator may create priority values (i.e. genes) that did not occur in the population before. Again, however, performing a mutation on an individual does not necessarily change the related schedule due to the redundancy discussed above.

5 Priority Rule based Genetic Algorithm

In this section, we describe a GA based on a priority rule representation. A similar GA has been proposed by Özdamar [24] for the multi-mode extension of the RCPSP. Like the priority value based one, also this GA employs the basic scheme and the selection operator variants that have been used for the permutation based GA. Next, the priority rule encoding and the related crossover and mutation operators are provided.

5.1 Individuals and fitness

In this GA variant, an individual I is given by a sequence of priority rules $pr_1^I, \dots, pr_J^I$ with

$$pr_i^I \in \{LFT, LST, MTS, MSLK, WRUP, GRPW\}$$

for each position $i = 1, \dots, J$. Each of these six priority rules has been suggested in the literature and shown to produce good schedules for the RCPSP, we refer to the study recently performed by Kolisch [14]. Table 1 contains a brief mathematical definition for each priority rule. Here, LFT_j denotes the latest finish time of activity j . It can be determined using traditional backward recursion from an upper bound on the project's makespan. Finally, with f_j we denote the earliest precedence and resource feasible finish time of activity j in the current partial schedule. Within the WRUP rule which has been developed by Ulusoy and Özdamar [28], we have employed the weights that performed best in their study.

Whereas Özdamar [24] employs a parallel scheme to transform an individual into a schedule, we use a serial scheme for the same reason as the one given in the previous section. First, we start the dummy source activity at time 0. Second, we compute the set of the eligible activities and use priority rule pr_1^I to select the eligible activity with the highest priority. The selected activity is started at the earliest precedence and resource feasible time. Repeatedly scheduling an unscheduled activity, we obtain a feasible active schedule. The fitness of an individual is again given by the makespan of the related schedule. Considering the project instance of Figure 1, the schedule related to the example genotype of Figure 2 (c) is again the one of Figure 3.

Each individual I of the initial generation is determined by randomly selecting one of the six priority rules for pr_i^I , $i = 1, \dots, J$.

Like the two previously described genetic representations, also the priority rule based encoding contains some redundancy. Consider again the example individual of Figure 2 (c). Replacing the first priority rule in the sequence (LFT) with e.g. the GRPW rule, we obtain a different genotype which is, however, also related to the schedule of Figure 3.

While the search spaces for the other two encodings always contain an optimal solution, this is not the case for the priority rule based representation. For the sake of shortness, we do not give a detailed counterexample, mentioning only that in some cases none of the employed priority rules may select an activity that must be scheduled next in order to obtain an optimal schedule from the current partial schedule. This drawback could be

Priority rule		formula	
LFT	latest finish time	min	LFT_j
LST	latest start time	min	$LFT_j - p_j$
MTS	most total successors	max	$ \bar{S}_j $
MSLK	minimum slack	min	$LFT_j - f_j$
WRUP	weighted resource utilisation and precedence	max	$0.7 S_j + 0.3 \sum_{k \in K} r_{jk} / R_k$
GRPW	greatest rank positional weight	max	$p_j + \sum_{i \in S_j} p_i$

Table 1: Good priority rules

overcome by adding a priority rule which allows any eligible activity to be selected, namely the random rule (RAND). However, we did not employ the RAND priority rule because it seems to be incompatible with the GA paradigm: Obviously, selecting an activity with the RAND rule may contribute to the fitness if, by chance, this activity continues the current partial schedule in an advantageous way. However, the RAND rule itself does not contain specific information that is worth to be inherited.

5.2 Crossover

We can employ standard crossover operators similar to those used for the priority value encoding. The definitions of the one-point, two-point, and uniform crossover for the priority rule representation are obtained from replacing pv_i^I by pr_i^I in the respective definitions for the priority value encoding.

5.3 Mutation

For the priority rule based encoding, the mutation operator is defined as follows: For each position $i = 1, \dots, J$ of an individual I , a new priority rule

$$pr_i^I \in \{\text{LFT}, \text{LST}, \text{MTS}, \text{MSLK}, \text{WRUP}, \text{GRPW}\}$$

is randomly drawn with a probability of p_{mutation} . Again, due to the redundancy described above, performing a mutation on a genotype does not necessarily change the related schedule.

6 Computational Results

6.1 Experimental Design

In this section we present the results of the computational studies concerning the genetic algorithms introduced in the previous sections. The experiments have been performed on a Pentium-based IBM-compatible personal computer with 133 MHz clock-pulse and 32 MB RAM. The procedures have been coded in ANSI C, compiled with the GNU C compiler and tested under Linux.

We used a set of standard test problems constructed by the project generator ProGen which has been developed by Kolisch et al. [18]. They are available in the project scheduling problem library PSPLIB from the University of Kiel. For detailed information the reader is referred to Kolisch and Sprecher [19].

In our study, we have used the problem sets containing instances with 30 and 60 non-dummy activities. The duration of a non-dummy activity varies between 1 and 10 periods. We have four renewable resources. For each problem size, a set consists of 480 instances which have been systematically generated by varying three parameters: network complexity, resource factor, and resource strength. The network complexity reflects the average number of immediate successors of an activity. The resource factor is a measure of the average number of resources requested per job. The resource strength describes the scarceness of the resource capacities. These parameters are known to have a big impact on the hardness of a project instance, cf. Kolisch et al. [18].

The set with 30 non-dummy activities currently is the hardest standard set of RCPSP-instances for which all optimal solutions are known, cf. Demeulemeester and Herroelen [6]. In what follows we report the average percentage deviation from the optimal makespan. However, for some of the instances with 60 activities, only heuristic solutions are known. In these cases, we give the average percentage deviation from the best lower and upper bounds as reported in the library PSPLIB at the time this research was performed. The lower bounds were computed by Baar et al. [1] and Heilmann and Schwindt [10]. The upper bounds were obtained by Kolisch and Drexel [17], Kohlmorgen et al. [12], and Baar et al. [1].

6.2 Configuration of the Genetic Algorithms

Crucial for the success of a GA is an appropriate choice of good genetic operators and adequate parameter settings, usually on the basis of computational experiments. We report the outcome of our study to determine the best GA configuration only for the suggested permutation based GA approach of Section 3 and for the instance set with 30 non-dummy activities, because the results for the other GAs and the other instance set are similar. For each instance, 1,000 schedules were computed.

Table 2 reports the average and the maximal deviation from the optimum, the percentage of instances for which an optimal solution was found, and the computation time in seconds for different combinations of alternative genetic operators. A mutation probability of 0.05, the two-point crossover strategy, and the ranking method for selection perform best. The two-point crossover operator appears to be capable of inheriting building blocks that contributed to the parents' fitness (for much larger projects, even more than two cuts may probably be advisable). In contrast, the uniform crossover operator (which yields good results for problems with a different structure such as e.g. the multidimensional knapsack problem, cf. Chu and Beasley [4]) does not seem to be well suited for sequencing problems. A randomized selection strategy seems to be advantageous only if a much larger number of individuals is considered.

p_{mutation}	crossover	selection	av. dev.	max. dev.	optimal	CPU-sec
0.01	two-point	ranking	0.64 %	9.7 %	77.5 %	0.54
0.05	two-point	ranking	0.54 %	7.9 %	81.5 %	0.54
0.10	two-point	ranking	0.56 %	8.6 %	80.4 %	0.55
0.05	one-point	ranking	0.65 %	9.7 %	77.5 %	0.54
0.05	two-point	ranking	0.54 %	7.9 %	81.5 %	0.54
0.05	uniform	ranking	0.66 %	8.6 %	79.6 %	0.76
0.05	two-point	ranking	0.54 %	7.9 %	81.5 %	0.54
0.05	two-point	proportional	0.62 %	7.7 %	78.5 %	0.60
0.05	two-point	2-tournament	0.63 %	9.3 %	79.0 %	0.54
0.05	two-point	3-tournament	0.59 %	7.3 %	80.9 %	0.54

Table 2: Alternative genetic operators — permutation based GA, 1,000 schedules, $J = 30$

Table 3 shows that a population size of 40 and 25 generations is the best parameter relationship when calculating 1,000 individuals (i.e. schedules). Finally, Table 4 shows that

it pays to use a random sampling method instead of a pure random procedure to determine the initial population.

<i>POP</i>	<i>GEN</i>	<i>av. dev.</i>	<i>max. dev.</i>	<i>optimal</i>	<i>CPU-sec</i>
50	20	0.56 %	11.1 %	80.8 %	0.54
40	25	0.54 %	7.9 %	81.5 %	0.54
20	50	0.79 %	9.2 %	76.5 %	0.54

Table 3: Impact of population size — permutation based GA, 1,000 schedules, $J = 30$

<i>Initial population</i>	<i>av. dev.</i>	<i>max. dev.</i>	<i>optimal</i>	<i>CPU-sec</i>
random sampling (LFT)	0.54 %	7.9 %	81.5 %	0.54
random	0.99 %	10.5 %	70.8 %	0.52

Table 4: Impact of initial population — permutation based GA, 1,000 schedules, $J = 30$

In the further computational studies, the three GAs make use of the best configuration determined here. That is, they apply a mutation probability of 0.05, a two-point crossover, the ranking selection, and the relationship of population size and number of generations given above. Recall, however, that the sampling procedure to determine the initial population is only employed in our permutation based GA. The other two approaches make use of a pure random procedure.

6.3 Comparison of the Approaches

In this subsection, we present the experimental results obtained from the comparison of the proposed GA of Section 3 with the GA approaches described in Sections 4 and 5.

The first numerical results to be presented are obtained from the project instance set with 30 activities where 1,000 individuals (i.e. schedules) are computed for each instance. Table 5 gives the average and the maximal deviation from the optimum, the percentage of instances for which an optimal solution was found, and the computation time in seconds. The permutation based encoding yields better results than the priority value based representation while both outperform the priority rule encoding. Also observe that the permutation based GA results in the lowest computation times. This is because we have to determine eligible activities and apply priority rules only for the first generation, when computing the schedule related to permutation based individuals.

Next, we examine the question whether the three representations are well suited for application in genetic algorithms. Clearly, we have two possibilities to use a given encoding: First, we can randomly generate a number of solutions, say 1,000, and choose the best one. Second, we can randomly generate only a few solutions, say 40, then apply the genetic operators over 25 generations, and choose the best among the resulting 1,000 solutions. Note that the former (random) procedure can be viewed as a GA with a population size of 1,000 and only one generation, that is, the initial one. Thus, the genetic operators are not applied there. The results of the random procedures for the encodings are listed in Table

<i>GA</i>	<i>av. dev.</i>	<i>max. dev.</i>	<i>optimal</i>	<i>CPU-sec</i>
permutation	0.54 %	7.9 %	81.5 %	0.54
priority value	1.03 %	10.8 %	70.6 %	0.64
priority rule	1.38 %	17.7 %	70.6 %	0.91

Table 5: Comparison of genetic algorithms — 1,000 schedules, $J = 30$

6. Comparing them to those of Table 5, we observe that the “real” GAs clearly outperform the corresponding random procedures if the permutation and priority value encodings are considered. However, for the priority rule encoding, there are only slight differences between randomly generating 1,000 schedules on the one hand and randomly generating only 40 and applying the genetic operators over 25 generations on the other hand. Hence we can deduce that the priority rule encoding is not well suited for genetic algorithms to solve the RCPSP while the other two representations are. Observe also that the initial population of the permutation based GA is computed using a random sampling method. That is, the random procedure for the permutation encoding listed in Table 6 is in fact a sampling algorithm. Comparing Tables 5 and 6, we see that it is better to determine only 40 schedules using the random sampling procedure and then apply the genetic operators over 25 generations, than to compute all 1,000 schedules with the random sampling method.

<i>Encoding</i>	<i>av. dev.</i>	<i>max. dev.</i>	<i>optimal</i>	<i>CPU-sec</i>
permutation	0.82 %	8.1 %	76.5 %	0.78
priority value	1.69 %	16.9 %	66.9 %	0.66
priority rule	1.41 %	16.2 %	69.8 %	0.93

Table 6: Comparison of encodings in random procedures — 1,000 schedules, $J = 30$

Up to this point, we have used a fixed number of schedules to be determined by the algorithms. However, the resulting computation times are different, see Table 5. Clearly, when using a heuristic procedure in practice, it is important to obtain good schedules within a reasonable amount of CPU time. Therefore we have additionally tested the three GA variants with time limits, but without bounding the number of schedules. As a further benchmark, we use a good sampling procedure known from the literature (cf. Kolisch [14]) which is based on the randomized LFT priority rule.

Table 7 displays the average deviations from the optimum obtained for the instances with $J = 30$ from four different time limits. The permutation based GA performs best for all time limits while the priority rule based GA yields the worst results. Note especially that the deviation of the permutation based GA is two times lower than that of the priority rule based GA for a time limit of 0.5 seconds while it is more than four times lower for 4 seconds. That is, the proposed GA is not only the best for small time limits, its superiority also further increases when the computation time is increased.

Next, we have performed the same experiment on the set of instances with 60 activities. As for some instances optimal solutions are currently unknown, we measure the deviations

<i>Algorithm</i>	<i>type</i>	<i>0.50 sec</i>	<i>1.00 sec</i>	<i>2.00 sec</i>	<i>4.00 sec</i>
GA	permutation	0.71 %	0.45 %	0.37 %	0.24 %
GA	priority value	1.16 %	0.88 %	0.69 %	0.54 %
GA	priority rule	1.51 %	1.33 %	1.21 %	1.13 %
sampling	LFT	1.00 %	0.77 %	0.66 %	0.55 %

Table 7: Average deviations w.r.t. time limit — $J = 30$

from the best known lower and upper bound here. The results can be found in Tables 8 and 9, respectively. Again, the permutation based GA performs best for all time limits. In contrast to the instances with $J = 30$, however, here the priority rule based GA outperforms the priority value GA and the sampling approach. This observation can be explained as follows: Selecting $J = 60$ instead of $J = 30$ results in a much larger search space. Within the same time limit, only a much smaller portion of the search space can be examined. Therefore, the strategy to combine several good priority rules corresponds to examining only potentially promising regions of the search space. However, the restriction to the regions identified by the priority rules is disadvantageous for smaller projects and/or higher computation times (or, of course, faster computers). Also keep in mind that the priority rule based GA hardly exploits the advantages of the genetic operators, cf. again Tables 5 and 6.

<i>Algorithm</i>	<i>type</i>	<i>0.50 sec</i>	<i>1.00 sec</i>	<i>2.00 sec</i>	<i>4.00 sec</i>
GA	permutation	5.30 %	4.73 %	4.37 %	4.16 %
GA	priority value	7.60 %	6.60 %	6.17 %	5.70 %
GA	priority rule	5.92 %	5.50 %	5.18 %	4.96 %
sampling	LFT	6.08 %	5.77 %	5.63 %	5.39 %

Table 8: Average deviations from best lower bound w.r.t. time limit — $J = 60$

<i>Algorithm</i>	<i>type</i>	<i>0.50 sec</i>	<i>1.00 sec</i>	<i>2.00 sec</i>	<i>4.00 sec</i>
GA	permutation	1.42 %	1.10 %	0.79 %	0.59 %
GA	priority value	3.61 %	2.88 %	2.24 %	1.79 %
GA	priority rule	2.07 %	1.80 %	1.52 %	1.34 %
sampling	LFT	2.27 %	1.99 %	1.82 %	1.62 %

Table 9: Average deviations from best upper bound w.r.t. time limit — $J = 60$

We remark here that most of the best known upper bounds for the hard instances were computed by Kohlmorgen et al. [12] with their parallel GA approach. Using a massively parallel computer with 16k processing units, they had 16,384 individuals per generation, while our GA could not evaluate more than 4,000 individuals altogether within 4 seconds when applied to the instances with 60 activities. Thus, it is not surprising that the best

known upper bounds are on the average 0.59 % better than our results (cf. Table 9). Nevertheless, it should be mentioned that the schedules for 3 of the 480 instances with 60 activities found by our GA (with a time limit of 4 seconds) were better than those reported in the library PSPLIB at the time this research was performed.

7 Conclusions

We have proposed a genetic algorithm (GA) for solving the classical resource-constrained project scheduling problem (RCPSP). The representation is based on a precedence feasible permutation of the set of the activities. The genotypes are transformed into schedules using a serial scheduling scheme. Among several alternative genetic operators for the permutation encoding, we chose a ranking selection strategy, a mutation probability of 0.05, and a two-point crossover operator which preserves precedence feasibility. The initial population was determined with a randomized priority rule method. In order to evaluate our approach, we have compared it to two GA concepts from the literature which make use of a priority value and a priority rule representation, respectively. As a further benchmark, a priority rule based random sampling procedure also known from the literature was tested. Our in-depth computational study revealed that our GA outperformed the other GAs as well as the sampling approach. This outcome suggests to apply our GA to real-world project scheduling problems. In fact, we have obtained promising results for a medical research project documented in [9].

References

- [1] BAAR, T., BRUCKER, P., AND S. KNUST (1997): Tabu-search algorithms for the resource-constrained project scheduling problem. Working Paper, University of Osnabrück, Germany.
- [2] BRUCKER, P., S. KNUST, A. SCHOO, AND O. THIELE (1996): A branch-and-bound algorithm for the resource-constrained project scheduling problem. *European Journal of Operational Research*, to appear.
- [3] CHO, J.H. AND Y.D. KIM (1997): A simulated annealing algorithm for resource-constrained project scheduling problems. *Journal of the Operational Research Society*, Vol. 48, pp. 736-744.
- [4] CHU, P.C. AND J.E. BEASLEY (1997): A genetic algorithm for the multidimensional knapsack problem. Working paper, The Management School, Imperial College, London.
- [5] DEMEULEMEESTER, E. AND W. HERROELEN (1992): A branch-and-bound procedure for the multiple resource-constrained project scheduling problem. *Management Science*, Vol. 38, pp. 1803-1818.
- [6] DEMEULEMEESTER, E. AND W. HERROELEN (1995): New benchmark results for the resource-constrained project scheduling problem. *Management Science*, to appear.
- [7] GOLDBERG, D.E. (1989): *Genetic algorithms in search, optimization, and machine learning*. Addison-Wesley, Reading, Massachusetts.
- [8] HARTMANN, S. (1997): Project scheduling with multiple modes: A genetic algorithm. *Manuskripte aus den Instituten für Betriebswirtschaftslehre*, No. 435, University of Kiel, Germany.

- [9] HARTMANN, S. (1997): Scheduling medical research experiments — An application of project scheduling methods. Manuskripte aus den Instituten für Betriebswirtschaftslehre, No. 452, University of Kiel, Germany.
- [10] HEILMANN, R. AND C. SCHWINDT (1997): Neue untere Schranken für das RCPSP/max. Paper presented on the Third Workshop über Projektplanung mit beschränkten Ressourcen, University of Kiel, Germany.
- [11] HOLLAND, H.J. (1975): Adaptation in natural and artificial systems. University of Michigan Press, Ann Arbor.
- [12] KOHLMORGEN, U., H. SCHMECK, AND K. HAASE (1996): Experiences with fine-grained parallel genetic algorithms. *Annals of Operations Research*, to appear.
- [13] KOLISCH, R. (1996): Efficient priority rules for the resource-constrained project scheduling problem. *Journal of Operations Management*, Vol. 14, pp. 179-192.
- [14] KOLISCH, R. (1996): Serial and parallel resource-constrained project scheduling methods revisited: Theory and computation. *European Journal of Operational Research*, Vol. 90, pp. 320-333.
- [15] KOLISCH, R. (1997): Resource allocation capabilities of commercial project management systems — Resource management boosts up the German Stock Exchange. Working paper, University of Kiel, Germany.
- [16] KOLISCH, R. AND K. HEMPEL (1996): Experimentelle Evaluation der Kapazitätsplanung von Projektmanagementsoftware. *Zeitschrift für betriebswirtschaftliche Forschung*, Vol. 48, pp. 999-1018.
- [17] KOLISCH, R. AND A. DREXL (1996): Adaptive search for solving hard project scheduling problems. *Naval Research Logistics*, Vol. 43, pp. 23-40.
- [18] KOLISCH, R., A. SPRECHER AND A. DREXL (1995): Characterization and generation of a general class of resource-constrained project scheduling problems. *Management Science*, Vol. 41, pp. 1693-1703.
- [19] KOLISCH, R. AND A. SPRECHER (1996): PSPLIB — A project scheduling problem library. *European Journal of Operational Research*, Vol. 96, pp. 205-216.
- [20] LEE, J.-K. AND Y.-D. KIM (1996): Search heuristics for resource-constrained project scheduling. *Journal of the Operational Research Society*, Vol. 47, pp. 678-689.
- [21] MICHALEWICZ, Z. (1995): Heuristic methods for evolutionary computation techniques. *Journal of Heuristics*, Vol. 1, pp. 177-206.
- [22] MINGOZZI, A., V. MANIEZZO, S. RICCIARDELLI AND L. BIANCO (1995): An exact algorithm for project scheduling with resource constraints based on a new mathematical formulation. *Management Science*, to appear.
- [23] MORI, M. AND C.C. TSENG (1997): A genetic algorithm for the multi-mode resource-constrained project scheduling problem. *European Journal of Operational Research*, Vol. 100, pp. 134-141.
- [24] ÖZDAMAR, L. (1996): A genetic algorithm approach to a general category project scheduling problem. Research Report, Marmara University, Istanbul.

- [25] REEVES, C.R. (1995): Genetic algorithms and combinatorial optimization. In: Rayward-Smith, V.J. (Ed.): Applications of modern heuristic methods. Alfred Waller Ltd., Henley-on-Thames.
- [26] SPRECHER, A. (1996): Solving the RCPSP efficiently at modest memory requirements. Manuskripte aus den Instituten für Betriebswirtschaftslehre, No. 425, University of Kiel, Germany.
- [27] SPRECHER, A., R. KOLISCH AND A. DREXL (1995): Semi-active, active and non-delay schedules for the resource-constrained project scheduling problem. European Journal of Operational Research, Vol. 80, pp. 94-102.
- [28] ULUSOY, G. AND L. ÖZDAMAR (1989): Heuristic performance and network/resource characteristics in resource-constrained project scheduling. Journal of the Operational Research Society, Vol. 40, pp. 1145-1152.