

Parallel and Sequential Testing of Design Alternatives

Christoph H. Loch

Christian Terwiesch

INSEAD

The Wharton School

Stefan Thomke

Harvard Business School

September 28, 1999

Abstract

An important managerial problem in product design is the extent to which testing activities are carried out in parallel or in series. Parallel testing has the advantage of proceeding more rapidly than serial testing but does not take advantage of the potential for learning between tests, thus resulting in a larger number of tests. We model this trade-off in form of a dynamic program and derive the optimal testing strategy (or mix of parallel and serial testing) that minimizes both the total cost and time of testing. We derive the optimal testing strategy as a function of testing cost, prior knowledge, and testing lead-time. Using information theory to measure the amount of learning between tests, we further show that in the case of imperfect testing (due to noise or simulated test conditions) the attractiveness of parallel strategies increases. Finally, we analyze the relationship between testing strategies and the structure of design hierarchy. We show that a key benefit of modular product architecture lies in the reduction of testing cost.

KEYWORDS: testing, prototyping, learning, optimal search, modularity

1 Introduction

Beginning with Simon (1969), a number of innovation researchers have studied the role of testing and experimentation in the research and development process (Simon, 1969; Allen, 1977; Wheelwright and Clark, 1992; Thomke, 1998; Iansiti, 1999). More specifically, Simon first proposed that one could “...think of the design process as involving, first, the generation of alternatives and, then, the testing of these alternatives against a whole array of requirements and constraints. There need not be merely a single generate-test cycle, but there can be a whole nested series of such cycles” (Simon, 1969, 1981, p. 149).

The notion of “design-test” cycles was later expanded by Clark and Fujimoto (1989) to “design-build-test” to emphasize the role of building prototypes in design, and to “design-build-run-analyze” by Thomke (1998) who identified the analysis of a test or an experiment to be an important part of the learning process in product design. These results echoed earlier empirical findings by Allen (1977; p. 60) who observed that research and development teams he studied spent on average 77.3% of their time on experimentation and analysis activities which were an important source of technical information for design engineers. Similarly, Cusumano and Selby (1996) later observed that Microsoft’s software testers accounted for 45% of its total development staff. Since testing is so central to product design, a growing number of researchers have started to study testing strategies, or to use Simon’s words once more, optimal structures for nesting a long series of design-test cycles (Cusumano and Selby, 1996; Thomke and Bell, 1999).

Integral to the structure of testing is the extent to which testing activities in design are carried out in parallel or in series. Parallel testing has the advantage of proceeding more rapidly than serial testing but does not take advantage of the potential for learning between tests - resulting in a larger number of tests to be carried out. As real-world testing strategies are combinations of serial and parallel strategies, managers and designers thus

face difficult choices in formulating an optimal policy for their firms. This is particularly important in a business context where new and rapidly advancing technologies are changing the economics of testing.

The purpose of this paper is to study the fundamental drivers of parallel and sequential testing strategies and develop optimal policies for research and development managers. We achieve this by formulating a model of testing that accounts for testing cost and lead-time, prior knowledge and learning between tests. We show formally under which conditions it is optimal to follow a more parallel or a more sequential approach. Moreover, using a hierarchical representation of design, we also show that there is a direct link between the optimal structure of testing activities and the structure of the underlying design itself; a relationship that was first explored by Alexander (1964) and later reinforced by Simon (1969, 1981).

Our analysis yields three important insights. *First*, the optimal mix of parallel and sequential testing depends on the *ratio* of the [financial] cost and [cost of] time of testing: More expensive tests make sequential testing more economical. In contrast, slower tests make parallel testing more attractive for development managers(see **Section 3**).

Second, imperfect tests reduce the amount of learning between testing sequential design alternatives. Using information theory to measure the amount of learning between tests, we show that such imperfect tests increase the attractiveness of parallel testing strategies(see **Section 4**).

Third, the structure of design hierarchy influences to what extent tests should be carried out in parallel or sequentially. We show that a modular product architecture can radically reduce testing cost compared to an integral architecture. We thus suggest a link between the extensive literature on design architecture and the more recent literature on testing (**Section 5**).

2 Parallel and Sequential Testing in Product Design

Design can be viewed as the creation of synthesized solutions in the form of products, processes or systems that satisfy perceived needs through the mapping between functional elements (FEs) and physical elements (PEs) of a product. Functional elements are the individual operations and transformations that contribute to the overall performance of the product. Physical elements are the parts, components, and sub-assemblies that implement the product's functions (Ulrich and Eppinger 1995, p. 131; see also Su 1990, p. 27).

To illustrate this view of product design, consider the following simple example. Assume that we are interested in designing the opening and closing mechanism of a door which has two FEs: the ability to *close* it (block it from randomly swinging open), with the possibility of opening from either side, and the ability to *lock* it (completely disallowing opening from one side or from both sides). The physical elements, or design alternatives, include various options of shape and material for the handle, the various barrels, and the lock (see Figure 1).

Insert Figure 1 about here

An integral characteristic of designing products with even moderate complexity is its *iterative nature*. As designers are engaged in problem-solving, they iteratively resolve *uncertainty* about which physical elements satisfy the perceived functional elements. We will refer to the resolution of this uncertainty as a test or a series of tests.

It is well-known that product developers generally do not expect to solve a design problem via a single iteration, and so often plan a series of design-test cycles, or experiments, to bring them to a satisfactory solution in an efficient manner (Allen, 1966; Simon, 1969; Smith and Eppinger, 1997; Thomke, 1998). When the identification of a solution to a design problem involves more than one such iteration, the information gained from a

previous test(s) may serve as an important input to the design of the next one. Design-test cycles which do incorporate learning derived from other cycles in a set are considered to have been conducted in series. Design-test cycles that are conducted according to an established plan that is not modified as a result of the finding from other experiments are considered to have been conducted in parallel.

For example, one might carry out a pre-planned “array” of design experiments, analyze the results of the entire array, and then carry out one or more additional verification experiments as it is the case in the field of formal “design of experiments (DOE)” methods (Montgomery 1991). The design-test cycles in the initial array are viewed as being carried out in parallel, while those in the second round are carried out in series with respect to that initial array. Such parallel strategies in R&D have been first suggested by researchers as far back as Nelson (1961) and Abernathy and Rosenbloom (1968), and more recently, by Thomke *et al.* (1998), and Dahan (1998).

Specifically, there are three important factors that influence optimal testing strategies: cost, learning between tests, and feedback time. First, a test’s *cost* typically involves the cost of using equipment, material, facilities, and engineering resources. This cost be very high, such as when a prototype of a new car is used in destructive crash testing, or it can be as low as a few dollars, such as when a chemical compound is used in pharmaceutical drug development and is made with the aid of combinatorial chemistry methods and tested via high-throughput screening technologies (Thomke *et al.* 1998). The cost to build a test prototype depends highly on the available technology and the degree of accuracy, or fidelity, that the underlying model is intended to have (Bohn 1987). For example, building the physical prototype used in automotive crash tests can cost hundreds of thousands of dollars whereas a lower-fidelity “virtual” prototype built inside a computer via mathematical modeling can be relatively inexpensive after the initial fixed investment in model building

has been made.

Second, the amount of *learning* that can be incorporated in subsequent tests is a function of several variables, including prior knowledge of the designer, the level of instrumentation and skill used to analyze test result, and, to a very significant extent, the topography of the “solution landscape” which the designer plans to explore when seeking a solution to her problem (Alchian, 1950; Kauffman and Levin, 1987; Baldwin and Clark, 1997a). In the absence of learning, there is no advantage of carrying out tests sequentially, other than meeting specific constraints that a firm may have (e.g. limited testing resources).

Third, the amount of learning is also a function of how timely feedback is received by the designer. It is well-known that misperceptions and delays in feedback from actions in complex environments can lead to suboptimal behavior and diminished learning. The same is true for noise which has shown to reduce the ability to improve operations (Bohn 1995). Thus, the *time* it takes to carry out a test and obtain results not only allows design work to proceed sooner but also influences the amount of learning between sequential tests.

3 A Model of Perfect Testing

We start our analysis by focussing on the optimal testing strategy in the design of one single physical element (PE). Consider for example the PE “locking mechanism” from Figure 1, for which there exist a number of design alternatives, depicted in Figure 2. Three different geometries of the locking barrel might fulfill the functional element (FE) “lock the door”. Based on her education and her previous work, the design engineer forms prior beliefs, e.g. “a cylinder is likely to be the best solution, however, we might also look at a rectangular prism as an alternative geometry”.

Insert Figure 2 about here

More formally, the engineer’s prior beliefs can be represented as a set of probabilities p_i defined over the alternatives $1..N$ where $p_i = \Pr\{\text{candidate } i \text{ is the best solution}\}$. In order to resolve the residual uncertainty, one geometry i is tested. Once the engineer can observe the result of the test, she gains additional information on whether or not this geometry is the best solution available. If a test resolves the uncertainty corresponding to a solution candidate completely, we refer to this test as a *perfect test* (imperfect testing will be analyzed in Section 4). Based on a test outcome, the designer can update her beliefs. If the tested candidate turns out to be the best solution, its probability gets updated to 1 and the other probabilities are renormalized accordingly. Otherwise, p_i is updated to 0. This updating mechanism represents learning in the model. It implies that a test reveals information on a solution candidate *relative* to the other candidates¹.

We assume that there is a fixed cost c per test as well as a fixed lead-time τ between the beginning of test-related activities and the observability of the newly generated information. The lead-time is important, as in presence of a delay, it can be beneficial to order several tests in parallel. Let c_τ be the cost of delay for the time-period of length τ . Testing thus “buys” information in form of updated probabilities at the price of $nc + c_\tau$, where n is the number of tests the engineer orders in one period.

For the special case $n = 1$, i.e. tests are done fully sequential, our testing problem can be seen as a search problem, similar to Weitzman (1979). In a result that is known as “Pandora’s rule”, Weitzman shows that if there are N “boxes” to be opened, box i offering a reward R with probability p_i , the box with the lowest “cost” $\frac{|A_i|c}{p(A_i)}$ should be opened first².

¹This corresponds to a situation where the design engineer can “tell the winner when she sees it”. As discussed above, this is one of many possible intermediate updates of the solution landscape.

²This review of Weitzman’s result has been adapted to correspond to our situation. In our problem, we consider less general rewards than in Weitzman’s Pandora’s rule (in our model, a candidate is either right or wrong, there is no generally distributed reward).

Here, $|A_i|$ is the number of objects in the box, c the search cost per object, and $p(A_i)$ the probability that the box contains the reward. Note that if all sets have equally many elements (in particular, if each solution candidate alone forms a set), this rule suggests to test the *most likely* candidate first.

However, Weitzman assumes that only one box can be opened at a time ($n = 1$), which ignores the aspect of testing lead-time. In most testing situations, the designer not only needs to decide *which* test to run next, but also *how many* tests should be run in parallel. On the one hand, parallel tests are attractive, as they resolve uncertainty in less time than sequential tests. In the extreme case where tests are carried out for every design alternative in parallel, the design problem is solved after one round of testing. On the other hand, parallel testing increases the number of tests as it fails to take advantage of the potential for learning from a test before running the next one. Learning is foregone as the designer commits to all tests at the same time. For a development manager, this creates an interesting trade-off between cost and time, which we will now explore further.

The described testing problem can be seen as a dynamic program, where the state of the system is the set S of remaining potential solution candidates with their probabilities. The decision to be made in each stage of the dynamic program is the set of states to be tested next, call it A . The immediate cost of this decision is $|A|c + c_\tau$, and the resulting state is the empty set with probability $p(A) = \sum_{i \in A} p_i$, and it is $S - A$ with probability $\sum_{i \in (S-A)} p_i$. A testing policy is optimal for a given set of solution candidates with attached probabilities p_i , if it minimizes the expected cost (testing and delay) of reaching the target state $S = \{\}$.

Theorem 1: To obtain the optimal testing policy, order the solution candidates in decreasing order of probability such that $p_i \geq p_{i+1}$. Assign the first candidates to set A_1 , the “batch” to be tested first, until its target probability specified in Equation (2) is reached.

Assign the next candidates to set A_2 to be tested next (if the solution is not found in A_1), and so on, until all N leaves are assigned to n sets $A_1, \dots, A_n$. The optimal number of sets³ is

$$n = \min \left\{ N; \max \left\{ 1, \left\lfloor \frac{1}{2} + \sqrt{\frac{1}{4} + \frac{2cN}{c_\tau}} \right\rfloor \right\} \right\}, \quad (1)$$

where $\lfloor \dots \rfloor$ denotes the integer part of a number. The sets are characterized by their probabilities $p(A_i) = \sum_{j \in A_i} p_j$:

$$p(A_i) = \frac{1}{n} + \frac{c_\tau}{cN} \left(\frac{n+1}{2} - i \right) = \frac{2(n-i)}{n(n-1)}. \quad (2)$$

It is interesting to note that the batch probabilities are described as a deviation from the average $1/n$: the first batches have a higher probability, the last batches a lower probability than the average. Note that this does *not* imply that the number of solution candidates in the first batches tested is also higher: if probabilities initially fall off steeply with i , the first batch tested may have a lower number of solution candidates than the second batch. If the total number of candidates N is very large, the difference in probability among the batches shrinks.

The policy in Theorem 1 behaves as we would intuitively expect. When the testing cost c is very large, the batches shrink to 1, $n = N$, and testing becomes purely sequential in order to minimize the probability that a given candidate must be tested. If c_τ approaches infinity, n approaches 1: testing becomes purely parallel in order to minimize time delay. When the total number of solution candidates N grows, the number of batches grows with $\sqrt{N}$. We describe this extreme behavior more precisely in the following corollary.

Corollary 1: If $\frac{1}{N} < \frac{c}{c_\tau} < \frac{N+1}{2}$, the optimal expected testing time is $\frac{n+1}{3}$, and the expected total testing cost is $\frac{c_\tau(n+1)(3n+2)}{12}$. If $\frac{c}{c_\tau} \leq \frac{1}{N}$, optimal testing is fully parallel

³The number of batches includes the last (n -th) set, which is empty. Thus, the *de facto* number of sets is $n - 1$.

($n = 1$), the testing time is 1, and the optimal total testing cost is $(c_\tau + Nc)$. If $\frac{c}{c_\tau} > \frac{N+1}{2}$, optimal testing is fully sequential, and the optimal total cost is $\sum_i ip_i(c + c_\tau)$. If all candidates are equally likely, this becomes $\frac{N+1}{2}(c + c_\tau)$.

In addition to defining the optimal testing policy, Theorem 1 provides an interesting structural insight concerning when to perform parallel search. Earlier studies have proposed that new testing technologies have significantly reduced the cost of testing, thus increasing the attractiveness of parallel strategies (e.g. Ward *et al.* 1995, Terwiesch *et al.* 1999, Thomke 1998). Our results clearly demonstrate this - as test cost decreases, the optimal batch size goes up. For the extreme case of $c = 0$, the above corollary prescribes a fully parallel search. This is precisely what happened in the pharmaceutical industry, when new technologies such as combinatorial chemistry and high-throughput screening reduced the cost of making and testing a chemical compound by orders of magnitude. Instead of synthesizing and evaluating, say, 5-10 chemical compounds per testing iteration, pharmaceutical firms now test for hundreds or thousands of compounds per test batch in the discovery and optimization of new drug molecules.

However, as the model shows, looking primarily at the cost benefits of new technologies ignores a second improvement opportunity. To fully understand the impact of new testing technologies on testing cost and search policy, one must consider that the results not only come at less cost, but that they also come in less *time*. In the automotive industry, for example, new prototyping technologies such as CAD based simulation or stereolithography have reduced the lead-time of a test to virtually zero. Thus, not only changes c , but so does c_τ .

Insert Figure 3 about here

If both parameters change *simultaneously*, the amount of parallel testing might go down or up. This interplay between testing cost and information turnaround times is illustrated

in Figure 3. The coordinates are speed ($\frac{1}{c_\tau}$) and cost effectiveness ($\frac{1}{c}$) of tests. The diagram in the lower left corner of the Figure represents testing economics with relatively low speed and cost effectiveness, resulting in some optimal combination of parallel and sequential testing as described in Theorem 1. Moving toward the lower right of the Figure corresponds to a reduction in testing cost, moving up to a reduction in testing time (or urgency). If a testing cost improvement outweighs a time improvement the test batches should grow, search becomes more parallel, as in the pharmaceutical example above.

If, in contrast, the dominant improvement is in the time dimension, the faster feedback time allows for learning between tests. The optimal search policy becomes “fast-sequential”. In this case, total testing cost *and* total testing time can decrease: total testing time because of shorter test lead-times and total testing cost because of “smarter” testing (based on the learning between tests, resulting in less wasted prototypes). Thus, in the evaluation of changing testing economics, a purely cost-based view may lead to an erroneous conclusion.

4 Imperfect Testing

Real-world testing is often carried out using simplified models of the test object (e.g. early prototypes) and the expected environment in which it will be used (e.g. laboratory environments). This results in imperfect tests. For example, aircraft designers often carry out tests on possible aircraft design alternatives using scale prototypes in a wind-tunnel - an apparatus with high wind velocities that partially simulate the aircraft’s intended operating environment. The value of using incomplete prototypes in testing is two-fold: to reduce investments in aspects of ‘reality’ that are irrelevant for the test, and to control out noise in order to simplify the analysis of test results. We model the effect of incomplete

tests and/or noise as *residual uncertainty* that remains after a design alternative has been tested (Thomke and Bell 1999). Such a test will be labeled as *imperfect*.

We assume that a test of design candidate i gives one of only two possible signals: $x = 1$ indicates “candidate i is the best design”, and $x = 0$ indicates “candidate i is not the best design”. An imperfect test of fidelity f is characterized by the conditional probabilities $p\{x = 1 \mid i = 1\} = 0.5(1 + f)$, and $p\{x = 1 \mid i = 0\} = 0.5(1 - f)$. The latter represents a “false positive,” and $\Pr\{\text{test} = 0 \mid i = 1\} = 0.5(1 - f)$ a “false negative”. To simplify exposition, we assume symmetry between the two errors. The test fidelity f captures the information provided by the test ($f = 0$: uninformative, $f = 1$: fully informative, perfect test). When $f < 1$, the probabilities can not be updated fully to 0 or 1, as we had assumed in Section 3.

This implies the following marginal probabilities of the signal from testing candidate i with fidelity f :

$$p\{x_i = 1\} = \frac{1}{2}[1 + f(2p_i - 1)]; \quad p\{x_i = 0\} = \frac{1}{2}[1 - f(2p_i - 1)]. \quad (3)$$

The posterior probabilities of all design candidates can be written as:

$$p\{i = 1 \mid x_i = 1\} = \frac{(1 + f)p_i}{1 + f(2p_i - 1)}; \quad p\{i = 1 \mid x_i = 0\} = \frac{(1 - f)p_i}{1 - f(2p_i - 1)}; \quad (4)$$

$$p\{j = 1 \mid x_i = 1\} = \frac{(1 - f)p_j}{1 + f(2p_i - 1)}; \quad p\{j = 1 \mid x_i = 0\} = \frac{(1 + f)p_j}{1 - f(2p_i - 1)} \quad (j \neq i). \quad (5)$$

If a test is perfect ($f = 1$), these posterior probabilities are the same as in the previous subsection. If a test is not perfect, it only *reduces the uncertainty* about a design alternative. It takes an infinite number of tests to reduce the uncertainty to zero (bring one p_k to 1). Therefore, the designer can only strive to reduce uncertainty of the design to a “sufficient confidence level $(1 - \alpha)$ ” in the design, where one $p_k \geq (1 - \alpha)$, and $\sum_{j \neq k} p_j \leq \alpha$. This is one of the reasons why a designer “satisfices”, as opposed to optimize, a product design (Simon 1969).

We first concentrate on a situation where only one alternative can be tested at once (sequential testing, Theorem 2a), turning to testing several alternatives in parallel afterward (Theorem 2b). The designer's problem is to find a testing sequence that reaches a sufficient confidence level at the minimum cost. As all information available to the designer is encapsulated in the system state $S = \mathbf{p} = \{p_1, \dots, p_N\}$ and the transition probabilities (4) and (5) depend only on S , we can formulate the problem as a dynamic program: At each test, pay an immediate cost of $(c + c_\tau)$ (for executing the test and for the time delay). Find a policy $\pi(\mathbf{p})$ that chooses a solution candidate $i \in \{1, \dots, N\}$ in order to minimize:

$$\begin{aligned} V(\mathbf{p}) = & (c + c_\tau) + \text{Min}_i \{ p\{x_i = 1\}V(p\{i = 1 \mid x_i = 1\}; p\{j = 1 \mid x_i = 1\} \ \forall j \neq i) \\ & + p\{x_i = 0\}V(p\{i = 1 \mid x_i = 0\}; p\{j = 1 \mid x_i = 0\} \ \forall j \neq i) \}, \end{aligned} \quad (6)$$

where $V(\mathbf{p}) = 0$ if and only if a design of sufficient confidence level has been found.

While we cannot write down the optimal testing cost for this problem, we can identify the optimal policy, showing that it has the same structure as for perfect tests.

Theorem 2a: *If testing is performed sequentially, that is, one design alternative at a time, it is optimal to always test the candidate with the largest p_i .*

Standard dynamic programming techniques cannot establish optimality of a myopic policy as stated in the theorem because the transition probabilities are state-dependent. Therefore, we use information theory as a tool to express the uncertainty reduction, or learning, offered by imperfect tests (Su 1990, Reinertsen 1997). This theory is based on Shannon (1948) and states that the *entropy* of a system indicates the amount of “choice” or uncertainty in that system. In particular, we define the *entropy* of the i^{th} design alternative and the entropy of the entire design problem, respectively, as

$$H_i = -p_i \log p_i, \quad (7)$$

$$H = \sum_i H_i. \quad (8)$$

The entropy captures knowledge about the alternative intuitively: It is maximal when $p_i = 1/2$, in which case $H_i = \log 2 = 1$ bit. That is, the uncertainty about design alternatives is maximal when all alternatives are equally likely to be the solution. $H_i = 0$ if $p_i = 0$ or if $p_i = 1$, that is, if it is known precisely whether the candidate leads to the solution or not.

The entropy H of the entire problem measures the uncertainty of the entire design. It is jointly concave in the p_i and maximal at $N \log N = N$ bits if all candidates are equally likely to be the best solution. $H = 0$ if and only if there is one candidate k with $p_k = 1$ (and thus, all other candidates are eliminated). Using the design problem's entropy, we can prove the Theorem (see Appendix).

Theorem 2a establishes that sequential testing with an imperfect fidelity f produces the same optimal order of alternatives to be tested – in order of decreasing probability. However, this order may change over the course of the testing procedure as the probabilities are updated.

Recall that a testing policy provides an assignment of testing candidates to time periods, or test batches (e.g. in the form of prototype sets). For perfect testing, this assignment could be done *ex-ante*, with the only exception being that the search should stop immediately after a positive test (see Theorem 1). The case of imperfect tests is harder, as the more complex updating of the probabilities makes an ex-ante assignment of tests to batches impossible. This is why Theorem 2a provides a dynamic policy.

We now relax the condition of sequentiality and allow the simultaneous testing of several design alternatives. We exclude multiple simultaneous tests of the *same* alternative⁴. We assume that the outcome of testing alternative i depends only on its own properties, but

⁴The situation does not correspond to, for example, consumer focus groups, where the same design alternative is shown to different consumers (which would increase the fidelity of the test).

not on any other alternative. The test outcomes are independent because of simultaneity – no learning takes place until after a test iteration has been completed. For parallel imperfect testing of this kind, we can prove the following result.

Theorem 2b: *Assume n different design alternatives are tested simultaneously as described above. Then it is optimal to always test the alternatives with the largest probabilities p_i . A higher number of parallel tests, n , reduces the entropy with diminishing returns, and there is an optimal number of parallel tests. If the test fidelity f decreases, the optimal number of parallel tests increases.*

Theorem 2b identifies an additional factor (to a lower ratio $\frac{c}{c_r}$) why parallel testing may be more economical. A lower testing fidelity diminishes the uncertainty reduction that can be gained from sequential tests. For any given size of a test batch, this increases the number of sequential rounds necessary to reach the target design confidence. As a result, a lower test fidelity, holding the testing cost c constant, increases the relative delay cost, and therefore increases the benefit of parallel testing. In the context of Figure 3, lower fidelity *de facto* reduces design speed, forcing more tests before finding the best alternative. Therefore, lower fidelity testing can lead to more parallelism.

5 Testing and the Structure of Design Hierarchy

A number of researchers have studied the role of design structure in the innovation process and have found it to matter significantly (Baldwin and Clark 1997a, Clark 1985, Marples 1961, Smith and Eppinger 1997, Simon 1969, Ulrich 1995). More specifically, it has been proposed that designs with smaller subsystems that can be designed and changed independently but function together as whole - a structure often referred to as modular - can have far-reaching implications for firm performance, including the management of prod-

uct development activities. This approach has been first explored by Alexander (1964) and was later reinforced by Simon (1969, 1981): “To design [such] a complex structure, one powerful technique is to discover viable ways of decomposing it into semi-independent components corresponding to its many function parts. The design of each component can then be carried out with some degree of independence of the design of other, since each will affect the others largely through its function and independently of the details of the mechanisms that accomplish the function” (Simon 1981, p. 148). In this section, we will explore the relationship between design structure and optimal testing.

A simple search model might capture the testing process related to one single physical element (PE) and a single functional element (FE), but in general, product design is concerned with more complex systems. The *design structure* links the product’s various FEs to its PEs. In the case of an uncoupled design, each FE is addressed by exactly one PE. Designs can also be coupled, in which case the mapping from FEs to PEs is more complex.

Insert Figure 4 about here

Consider the two different door designs illustrated in Figure 4. The design on the left of Figure 4 is *uncoupled*, that is, each FE is addressed by exactly one physically separate component. Closing is performed by a handle that moves a blocking barrel (which inserts into the door frame), and locking is carried out by turning a key that moves a second barrel. If the design is uncoupled, each FE is fulfilled by one PE, and each PE contributes to one FE. We call this separation of FEs *functional independence* of the design⁵. Designs that are functionally independent, are also called *modular* (Ulrich and Eppinger 1995).

⁵In addition to functional dependencies, elements can also dependent on each other because of their physical attributes which we will refer to as technical dependence. The interdependence between PEs can be captured in the *design structure matrix* (DSM) (Steward 1981, Eppinger et al. 1994).

The design on the right in Figure 4 is *coupled*. Closing is implemented by a doorknob, the turning of which moves a blocking barrel. Locking is enacted by a button in the center of the doorknob that blocks the doorknob from turning. The locking function uses *both* physical components, in particular, the same rod moving the barrel when opening/closing the door is blocked from moving when locking the door.

The architecture of the product has a fundamental influence on the testing process. In the case of functional independence between closing and locking the door, the corresponding subsystems (PEs) can be tested independently. If there are three candidates for the barrel (Figure 4) and two candidates for the lock, a total of $3+2=5$ tests would cover the total search space. If, however, the closing and locking are coupled, testing requires a specification of both PEs, closing barrel *and* locking barrel. If the outcome of the test is negative (FEs were not fulfilled), learning from the failure is more complex. E.g., if the closing FE was fulfilled, but not the locking FE, the engineer can not infer whether she should just change the locking barrel, or also the closing barrel. An exhaustive search requires $3*2=6$ tests⁶.

An intermediate case between coupled design and uncoupled design results, if the PEs contributing to the first FE can be determined without specifying the PEs contributing to the second FE, but not vice versa. In this case, we speak of *sequential dependence*, and

⁶Simon (1969) illustrates this point very nicely with the following example, which was originally supplied by W. Ross Ashby. "Suppose that the task is to open a safe whose lock has 10 dials, each with 100 possible settings, numbered from 0 to 99. How long will it take to open the safe by a blind trial-and-error search for the correct setting? Since there are 100^{10} possible settings, we may expect to examine about one half of these, on average, before finding the correct one [...]. Suppose, however, that the safe is defective, so that a click can be heard when anyone dial is turned to the correct setting. Now each dial can be adjusted independently and does not need to be touched again while the others are being set. The total number of settings that have to be tried is only 10×50 , or 500."

it is possible to test the first PE/FE before addressing the second.

We see that functional and technical structure influences testing in two ways. First, it influences the number of tests required for an exhaustive search ($3*2$ vs $3+2$ in the door example). Second, it influences the timing of the tests. If the design is uncoupled, tests can be done in parallel (without any additional cost). In the case of sequential dependence, parallel testing is possible, but only up to a certain level. Coupled designs however, cause the search space to grow exponentially without opportunities for parallel testing (other than the parallel testing where the designer precommits to several prototypes at once). The resulting effect of product architecture on testing cost is analyzed in Theorem 3 below.

For simplicity of exposition, assume a symmetric situation where each PE has N solution alternatives of equal probability, and there is one PE for each of M functional requirements. We consider the three generic architectures independent (modular), sequentially dependent (any two PEs have an upstream-downstream relationship), or integrated (each PE impacts all other PEs). Clearly, most complex systems include aspects of all three of these categories, but in the interest of a clear comparison, it is most useful to analyze them as three distinct types along a spectrum of structural possibilities.

Theorem 3: Suppose a design has M PEs with N equally likely solution candidates each, and a test costs c and takes one time unit costing c_τ . Then the expected testing costs for the three architectures are (where $n(N) = 1/2 + \sqrt{1/4 + 2cN/c_\tau}$ from Theorem 1):

	Parallel, $n(N) = 1$ $(\frac{c}{c_\tau} \leq \frac{1}{N})$	Intermediate $(\frac{1}{N} < \frac{c}{c_\tau} < \frac{N+1}{2})$	Sequential, $n(N) = N$ $(\frac{N+1}{2} \leq \frac{c}{c_\tau})$
$\mathbf{C}_{\text{mod}} \leq$	$c_\tau + NMc$	$Mc_\tau \left[\frac{n(N)}{M+1} + \frac{(n(N)+1)(n(N)-2/3)}{12} \right]$	$\frac{MN}{M+1} (c_\tau + c)$
$\mathbf{C}_{\text{sequ}} \leq$	$Mc_\tau + NMc$	$Mc_\tau \frac{(n(N)+1)(n(N)-2/3)}{12}$	$\frac{N+1}{2} M (c_\tau + c)$
$\mathbf{C}_{\text{int}} =$	$c_\tau + N^M c$ (if $\frac{c}{c_\tau} \leq \frac{1}{N^M}$)	$\frac{c_\tau}{12} [n(N^M) + 1][n(N^M) - \frac{2}{3}]$ (if $\frac{1}{N^M} < \frac{c}{c_\tau} < \frac{N^M+1}{2}$)	$\frac{N^M+1}{2} (c_\tau + c)$ (if $\frac{N^M+1}{2} \leq \frac{c}{c_\tau}$)

Corollary 2: The testing costs are affected by the product architecture as follows: $C_{mod} < C_{sequ} \leq C_{int}$.

The Theorem shows that in a modular architecture, the expected testing costs grow *sub-linearly* with the number of PEs; the costs grow *linearly* in a sequentially dependent architecture, and they grow *exponentially* in an integrated architecture. Thus, the Theorem explains that *testing effort* contributes to the benefits of a modular product architecture, simplifying the development process and often leading to lower total development time and cost. Figure 5 summarizes the connection between architecture and testing.

Insert Figure 5 about here

The results of Theorem 3 are consistent with similar propositions in the literature. Ulrich (1995) noted that for modular architectures, the design of each module can proceed almost independently and in parallel. System-level product testing would be limited to detecting unanticipated interactions, or areas where the system is not perfectly modular. The result of our analysis shows this to be true if modularity can be established in the functional and physical domains and if there is a direct one-to-one mapping between functional and physical elements (FEs and PEs). In such an extreme case of functional and technical modularity, there is no need for system-level testing. However, if there is at least one FE that is impacted by all PEs, the benefits of modularity are substantially reduced. In fact, all design alternatives and their impact on this FE would have to be considered for testing - a number that would increase very rapidly as Theorem 3 shows. If designers know little about functional [customer] elements and their interactions - not an unusual real-world dilemma - the value of modularity in testing quickly disappears. Indeed, as Baldwin and Clark (1997b) have shown, the presence of many modules can lead to a combinatorial explosion of testing and experimentation if the system-level impact on markets (or, in our definition, on the functional user domain) is highly uncertain.

6 Conclusion

In this paper, we have shown that the extent to which testing activities are carried out in parallel and series can have a significant impact on design performance. Parallel testing has the advantage of proceeding more rapidly than serial testing but does not take advantage of the potential for learning between tests, thus resulting in a larger number of tests. We model this trade-off in form of a dynamic program and derive the optimal testing strategy (or mix of parallel and serial testing) that minimizes both, the total cost of testing and time. More specifically, our paper shows three results.

First, the optimal mix of parallel and sequential testing depends on the *ratio* of testing cost and time: More expensive tests make sequential testing more economical. In contrast, slower tests or an increasing opportunity cost of time make parallel testing more attractive for development managers.

Second, imperfect tests reduce the amount of learning between testing sequential design alternatives and thus increase the attractiveness of parallel testing strategies. This is particularly important for managers who consider switching to early and less complete prototypes and/or the use of less controlled test environments.

Third, the design structure influences to what extent tests should be carried out in parallel or sequentially. We show that an important benefit of a modular product architecture comes from reduced testing cost, because it allows parallel testing without an increase in the number of test combinations. Thus, architecture can be an important lever for decreasing test cost.

As part of our testing model, we were also able to extend an important search model developed by Weitzman (1979). Whereas Weitzman studied sequential search, we included the option of carrying out search (or, in our case, testing) in parallel. We derived policies

that would not only prescribe an optimal sequence of search but also inform decision-makers about the degree to which such searches should be carried out in parallel. We expanded our analysis to include imperfect testing and, using principles from information theory to model uncertainty reduction, examined the impact of reduced learning on testing. Our last theorem also confirmed that there is an important relationship between two streams of research (design structure and testing) which we tried to establish more formally.

To conclude this paper, we propose three promising directions for further research that build directly on the findings presented here. First, it has been empirically observed that iterative testing can not only influence development cost and time but also the quality of the design solution. It has been found that less costly and faster iterations through advanced technologies such as computer simulation can actually result in more experimentation, leading to novel solution concepts that could not be reasonably tested for with slower and more costly technologies. In the present paper, we have focused on the cost and time aspects, holding design solution quality constant. But it is possible to make N , the number of design alternatives tested, an explicit decision variable in our model. Expanding the search space will increase the testing costs, but also improve the design quality (possibly with diminishing returns). This future work relates our current model to the literature on set-based product development (e.g., Ward 1995).

Second, in the case of sequential dependence, it might be beneficial to start the testing of the second module before the testing for the first module has been finalized, i.e. to overlap the two testing processes, in the spirit of concurrent engineering (Loch and Terwiesch 1998). Finding the optimal level of overlap between tests is thus a second opportunity for further research.

Third, we have shown learning between tests to be the primary advantage of sequential testing. In this paper, we have modeled the consequences of learning (uncertainty reduc-

tion through probability updating) but have not explicitly taken advantage of what is known about the different factors that influence learning. For example, a solution landscape represents the arena that the designers search to identify a solution to their problem. The probability of finding a solution increases as one ascends the “hills” in the landscape, and so the designer’s goal is to devise a series of tests that will enable them to identify and explore those hills in an efficient manner. The amount that can be learned between tests then relates directly to the topography of the landscape. Very little can be learned by a designer about the direction of her search if, for example, the solution landscape is absolutely flat for all combinations except the correct one. In contrast, suppose that the solution landscape is a hill with only a single peak and sides that extend to all edges of the landscape⁷. In such a case a strategy of serial testing may be the most efficient choice, because the information gained from each step taken is so useful in guiding the direction of the next trial step that the correct solution is often found after only a few trials.

Certainly, knowledge about the topology of solution landscapes will make sequential testing more attractive to designers. It is thus not surprising that well-studied engineering design problems tend to follow more sequential plans than, say, the early search for drug candidates in a relatively unknown solution space such as Alzheimer’s disease, even after the cost and time of each test is accounted for. Some of these factors that influence learning can be included more explicitly in our model of parallel and sequential testing and thus provide further leverage in the formulation of optimal testing strategies for superior product development performance.

References

Abernathy, W. and R. Rosenbloom. “Parallel and Sequential R&D strategies: application of a

⁷This is the shape, for example, of the value landscape in the children’s’ game in which a child is guided to a particular spot via feedback from other children who say ”warmer” each time a step is taken towards that spot.

- simple model,” *IEEE Transactions on Engineering Management*, 15 (1), March 1968.
- Alchian, A. “Uncertainty, Evolution and Economic Theory,” *Journal of Political Economy*, 58 (3), 1950.
- Alexander, C. *Notes on the Synthesis of Form*. Cambridge: Harvard University Press 1964.
- Allen, T. J. “Studies of the Problem-Solving Process in Engineering Design”, *IEEE Transactions on Engineering Management*, EM-13, no.2: 72-83, 1966.
- Allen, T.J. *Managing the Flow of Technology*, MIT Press, Cambridge, Massachusetts 1977.
- Baldwin, C. and K. Clark. “Design Options and Design Evolution,” Harvard Business School Working Paper 97-038, 1997a.
- Baldwin, C. and K. Clark. “The Value of Modularity: Splitting and Substitution,” Harvard Business School Working Paper 97-039, 1997b.
- Bertsekas, D. P. *Dynamic Programming and Optimal Control*, Belmont, Mass.: Athena Scientific 1995.
- Bohn, Roger. “Learning by Experimentation in Manufacturing,” Working Paper No. 88-001, Harvard Business School, 1987.
- Bohn, R. E., “Noise and Learning in Semiconductor Manufacturing” , *Management Science*, 1995, 41(1): 31-42
- Clark, K. B. “The Interaction of Design Hierarchies and Market Concepts in Technological Evolution.” *Research Policy* 14, 1985, 235 - 251.
- Clark, K. B., and T. Fujimoto. “Lead Time in Automobile Development: Explaining the Japanese Advantage, *Journal of Technology and Engineering Management* 6, 1989, 25-58.
- Cusumano, M., and R. Selby. *Microsoft Secrets*, New York: The Free Press 1995.
- Dahan, E. *Parallel and Sequential Prototyping in Product Development*, Unpublished Ph.D. Dis-

sertation, Stanford University, 1998.

Ha, A. Y., and E. L. Porteus, "Optimal Timing of Reviews in Concurrent Design for Manufacturability." *Management Science*, 41, 1431-1447

Iansiti, M. "How the Incumbent Can Win: Managing Technological Transitions in the Semiconductor Industry," Harvard Business School Working Paper 99-093, 1999, forthcoming in *Management Science*.

Kauffman, S. and S. Levin. "Towards a General Theory of Adaptive Walks on Rugged Landscapes," *Journal of Theoretical Biology* 128, 1987, 11-45.

Loch, Christoph H., Christian Terwiesch, "Communication and Uncertainty in Concurrent Engineering", *Management Science*, Vol. 44, No. 8, 1998

Marples. D. L. "The Decisions of Engineering Design." *IRE Transactions on Engineering Management* 1961, 55 - 71.

Montgomery, D. *Design and Analysis of Experiments*. New York: Wiley, 1991.

Nelson, R. "Uncertainty, Learning, and the Economics of Parallel Research and Development Efforts," *Review of Economics and Statistics* 43, 1961, 351-364.

Reinertsen, D., *Managing the Design Factory*. New York: The Free Press 1997.

Shannon, C. E. "A Mathematical Theory of Communication." *Bell Systems Technical Journal* 27, 1948, 379 - 423 and 623 - 656.

Simon, H. A. *The Sciences of the Artificial*. Cambridge, MA: MIT Press 1969.

Smith, R. P., and S. D. Eppinger. "A Predictive Model of Sequential Iteration in Engineering Design," *Management Science* 43, 1997, 1104 - 1120.

Steward, D. V. *Systems Analysis and Management: Structure, Strategy, and Design*. New York: Petrocelli Books, 1981.

Su, N. P. *The Principles of Design*, Oxford: Oxford University Press 1990.

Terwiesch, C., C. H. Loch, and A. De Meyer. “Exchanging Preliminary Information in Concurrent Development Processes.” Wharton/INSEAD Working Paper 1999.

Thomke, S. “Managing Experimentation in the Design of New Products,” *Management Science* 44, 1998, 743-762.

Thomke, S. and Bell, D. “Optimal Testing in Product Development,” Harvard Business School Working Paper No. 99-053, 1999.

Thomke, S., E. von Hippel and R. Franke, “Modes of Experimentation: An Innovation Process - and Competitive - Variable,” *Research Policy* 1998, 315-332.

Ulrich, K. “The Role of Product Architecture in the Manufacturing Firm.” *Research Policy* 24, 1995, 419 - 440.

Ulrich, K. and S. Eppinger. *Product Design and Development*, McGraw-Hill, New York, 1994.

Ward, A., J. Liker, J. Cristiano and D. Sobek. “The Second Toyota Paradox: How Delaying Decisions Can Make Better Cars Faster,” *Sloan Management Review*, Spring 1995, 43 - 61.

Weitzman, M. L. “Optimal Search for the Best Alternative.” *Econometrica* 47, 1979, 641 - 654.

Wheelwright, S. C., and K. B. Clark. *Revolutionizing Product Development*. New York: The Free Press 1992.

Appendix

Proof of Theorem 1. If the solution is not found within the set A_i , the next set must be tested, which happens with probability $\frac{1 - \sum_{j=1}^i p(A_j)}{1 - \sum_{j=1}^{i-1} p(A_j)}$ (the denominator updates the probabilities of the remainign sets to sum to 1). Thus, we can write the total expected search cost as

$$\begin{aligned} EC = & | A_1 | c + c_\tau + (1 - p(A_1)) [| A_2 | c + c_\tau + \frac{1 - p(A_1) - p(A_2)}{1 - p(A_1)} [| A_3 | c + c_\tau \\ & + \frac{1 - p(A_1) - p(A_2) - p(A_3)}{1 - p(A_1) - p(A_2)} [\dots]]] \end{aligned}$$

$$= \sum_{i=1}^n [(|A_i| + c + c_\tau) \sum_{j=i}^n p(A_j)]. \quad (9)$$

The fact that the design alternatives should be assigned in decreasing order of probability follows from an exchange argument: Assume that there are two alternatives $j \in A_i$ and $k \in A_{i+1}$ with $p_k > p_j$. Exchange the two (test k before j). The resulting change in total expected cost is, from (9), $(c_\tau + |A_{i+1}|)(p_j - p_k) < 0$. Thus, the candidates should be assigned as stated.

To simplify exposition, assume from now on that N is sufficiently large and the p_i small to approximate them by a continuous distribution function F . Now we transform the space, considering instead of the set sizes $|A_i|$ their probabilities $a_i = F(A_i) - F(A_{i-1})$, with $\sum_i a_i = 1$. The set sizes a_i correspond to fractions of N . In the transformed space, the solution candidates have a uniform probability density of 1, and the testing cost becomes Nc because the number of candidates has been compressed from N to 1. We can now state the objective function to be minimized (where we leave out the constraint that $n \leq N$ as it can be easily incorporated at the end):

$$\text{Min}_{n, a_i} \quad \sum_{i=1}^n (c_\tau + a_i Nc) \left(1 - \sum_{j=1}^{i-1} a_j\right) \quad (10)$$

$$\text{subject to} \quad \sum_j a_j = 1; \quad a_i \geq 0 \quad \forall i. \quad (11)$$

The Lagrangian of this objective function is $L = c_\tau \sum_i i a_i + Nc \sum_i a_i (1 - \sum_{j=1}^{i-1} a_j) - \lambda (1 - \sum_i a_i) - \sum_i \mu_i a_i$. The optimality conditions for the Lagrangian are $a_i \frac{\partial L}{\partial a_i} = 0 \forall i$, $\frac{\partial L}{\partial \lambda} = 0$, and $\mu_i \frac{\partial L}{\partial \mu_i} = 0 \forall i$. These, in turn, yield the condition

$$c_\tau k + Nc a_k + \lambda = 0 \quad \text{for all } k \text{ such that } a_k > 0. \quad (12)$$

Condition (12), first, implies that the second order condition is fulfilled (differentiating it with respect to a_k gives $Nc > 0$, so the solution found is a cost minimum). Second, (12) implies that the sets a_k are decreasing in size over k , so the first n^* sets are non-empty, and then no more candidates are assigned. Adding Equation (12) over all k and using

the fact that $\sum_k a_k = 1$ allows determining λ , and substituting in λ yields the optimal set probability (2). Finally, when the set probabilities are known, we can use the fact that $a_{n^*} > 0$ and $a_{(n+1)^*} \leq 0$ to calculate the optimal number of sets described in Equation (1). If $n^* \geq N$, then every solution candidate is tested by itself, which yields the largest number of sets possible.

Proof of Theorem 2a: We prove the theorem in three steps.

Step 1: Equivalence to entropy reduction. As the immediate reward $-(c + c_\tau)$ is constant, the problem is to minimize the expected number of steps to go from the initial state to $V = 0$ (Bertsekas 1995, 300). Moreover, $H(\mathbf{p})$ is unique given $\mathbf{p}$, and there is a unique state (up to permutations of the alternatives) producing the entropy $H_0 \equiv H(1 - \alpha, \frac{\alpha}{N}, \dots, \frac{\alpha}{N})$, namely the same state that yields $V = 0$. Thus, getting from $V(\mathbf{p})$ to $V = 0$ is equivalent to getting from $H(\mathbf{p})$ to H_0 .

Step 2: One-step entropy reduction. After testing design alternative i , we can write the posterior entropy as (where “ $x_i = a$ ” is abbreviated as “ a ”):

$$\begin{aligned}
H_{post} &= -p\{x_i = 1\}[p\{i = 1 \mid 1\} \log(p\{i = 1 \mid 1\}) + \sum_{j \neq i} p\{j = 1 \mid 1\} \log(p\{j = 1 \mid 1\})] \\
&\quad -p\{x_i = 0\}[p\{i = 1 \mid 0\} \log(p\{i = 1 \mid 0\}) + \sum_{j \neq i} p\{j = 1 \mid 0\} \log(p\{j = 1 \mid 0\})] \\
&= -\frac{1}{2} \left\{ (1+f)p_i \log\left(\frac{(1+f)p_i}{1+f(2p_i-1)}\right) + \sum_{j \neq i} (1+f)p_j \log\left(\frac{(1+f)p_j}{1+f(2p_i-1)}\right) \right. \\
&\quad \left. + (1-f)p_i \log\left(\frac{(1-f)p_i}{1-f(2p_i-1)}\right) + \sum_{j \neq i} (1-f)p_j \log\left(\frac{(1-f)p_j}{1-f(2p_i-1)}\right) \right\} \quad (13)
\end{aligned}$$

$$\begin{aligned}
&= \frac{1}{2} \left\{ \sum_{j=1}^N [(1+f)p_j \log((1+f)p_j) + (1-f)p_j \log((1-f)p_j)] \right. \\
&\quad \left. + [1+f(2p_i-1)] \log[1+f(2p_i-1)] + [1-f(2p_i-1)] \log[1-f(2p_i-1)] \right\}. \quad (14)
\end{aligned}$$

(13) tells us that $H_{post} < H(\mathbf{p})$ no matter which candidate is tested. This is because H_{post} is a symmetric-spread-out combination of the summands of H , which is smaller because H is jointly concave.

(14) tells us that H_{post} is minimal if the candidate i is tested which has a probability closest

to $1/2$: the left summand is independent of which candidate i is tested (and it is smaller than H , again because of concavity). The right summand of (14) is a *convex* function of p_i of the form $(1 + \epsilon) \log(1 + \epsilon) + (1 - \epsilon) \log(1 - \epsilon)$, where $1 \geq \epsilon = f | 2p_i - 1 | \geq 0$, and moreover, the right-hand summand is zero if $\epsilon = 0$ or $p_i = 1/2$. Therefore, the right summand is smallest if p_i is closest to $1/2$.

To conclude step 2, we observe that p_i being closest to $1/2$ is equivalent to p_i being the largest: If all $p_j \leq 1/2$, this is true trivially. If one $p_k > 1/2$, then $p_k - 1/2 = 1/2 - \sum_{j \neq k} p_j$ which implies that p_k is closer to $1/2$ than all the p_j . This shows step 2. We now need to show step 3, that this is also the optimal stationary policy.

Step 3: Optimal stationary policy. We examine the expected *two-step* entropy reduction assuming that candidates i and then k are tested, while $j \neq i, k$ refers to all remaining candidates. Four cases result, of the test signals being $(1, 1), (1, 0), (0, 1), (0, 0)$. Because of renormalization, the updated probabilities are arithmetically messy, and we give them only for the case $(1, 1)$ (the others are left to the reader or can be obtained from the authors):

$$\begin{aligned} p\{x_i = x_k = 1\} &= \frac{1 + 2f(1 - f)(p_i + p_k) - f(2 - f)}{4}; & p_i(post) &= \frac{(1 - f)(1 + f)p_i}{4p\{x_i = x_k = 1\}}; \\ p_k(post) &= \frac{(1 - f)(1 + f)p_k}{4p\{x_i = x_k = 1\}}; & p_j(post) &= \frac{(1 - f)(1 - f)p_j}{4p\{x_i = x_k = 1\}}. \end{aligned} \quad (15)$$

The resulting two-step posterior entropy becomes

$$\begin{aligned} H_{post^2} &= -\frac{1}{4} \{ \bar{H} - [1 + 2f(p_i + p_k) - f] \log[1 + 2f(1 - f)(p_i + p_k) - f(2 - f)] \\ &\quad - [1 + 2f(1 + f)p_i - 2f(1 - f)p_k - f^2] \log[1 + 2f(1 + f)p_i - 2f(1 - f)p_k - f^2] \\ &\quad - [1 - 2f(1 - f)p_i + 2f(1 + f)p_k - f^2] \log[1 - 2f(1 - f)p_i + 2f(1 + f)p_k - f^2] \\ &\quad - [1 - 2f(p_i + p_k) + f] \log[1 - 2f(1 + f)(p_i + p_k) + f(2 + f)] \}; \end{aligned} \quad (16)$$

where $\bar{H} = \sum_m \{(1 - f)^2 p_m \log[(1 - f)^2 p_m] + 2(1 - f)(1 + f)p_m \log[(1 - f)(1 + f)p_m] + (1 + f)^2 p_m \log[(1 + f)^2 p_m]\}.$

Inspection shows that H_{post^2} is the same when the order of testing i and k is exchanged. By induction, this implies that any order of testing a given collection of candidates gives in expectation the same posterior entropy. The result from step 2 that testing the largest p_i in the first round yields the largest entropy reduction, together with the result of step

3 that the order of testing a given collection does not matter, implies that it is optimal to test the largest p_i in all rounds. This completes the proof of the theorem. $\square$

Proof of Theorem 2b: When we test design alternatives $i = 1, \dots, n$ in parallel, our independence assumption implies that test outcome x_i is determined by (3), no matter what the other alternatives and tests are. Therefore, the conditional probability of $(x_1, \dots, x_n \mid i = 1, \dots, n) = (1/2^n)(1+f)^{n_r}(1-f)^{n-n_r}$, where n_r is the number of tests that give the “right” signal ($x_i = 1$ iff $i = 1$), and $n - n_r$ the number of tests that give the wrong signal. Recall that it is impossible that more than one of the alternatives is in fact the right one.

Now consider an arbitrary profile of test signals $\mathbf{x} = (x_1, \dots, x_n)$. Denote by $\mathcal{K}$ the subset of the n tested candidates for which the test signal is positive: $x_k = 1$ for $k \in \mathcal{K}$, and write the size of $\mathcal{K}$ as K . The marginal probability of the profile $\mathbf{x}$ is:

$$p(\mathbf{x}) = \frac{(1-f)^K(1+f)^{n-K}}{2^n} R, \quad (17)$$

where $R = [1 + \frac{2f \sum_{k \in \mathcal{K}} p_k}{1-f} - \frac{2f \sum_{m \notin \mathcal{K}} p_m}{1+f}]$. It represents a probability update after the tested candidates have changed their probabilities. The posterior probabilities follow.

$$x_l \text{ tested:} \quad x_l = 1 : p\{l = 1 \mid \mathbf{x}\} = \frac{(1+f)p_l}{(1-f)R}; \quad x_l = 0 : p\{l = 1 \mid \mathbf{x}\} = \frac{(1-f)p_l}{(1+f)R} \quad (18)$$

$$x_j \text{ not tested:} \quad p\{j = 1 \mid \mathbf{x}\} = \frac{p_j}{R}. \quad (19)$$

The posterior entropy of a tested candidate i must be taken over *all* possible signal profiles $\mathbf{x}$ with any set $\mathcal{K}$ of K positive tests, for any number K of positive tests. With j , we refer to a not tested candidate. Note that there are $\binom{n}{K}$ different sets $\mathcal{K}$ with K positive signals, and $\sum_{K=1}^n \binom{n}{K} = 2^n$. Thus, the denominator in the posterior entropy below represent a normalization:

$$\begin{aligned} H_n(i) = & - \sum_{K=0}^n \frac{(1-f)^K(1+f)^{n-K}}{2^n} \left\{ \sum_{\mathcal{K}: i \in \mathcal{K}} \frac{1+f}{1-f} p_i \log\left[\frac{1+f}{1-f} p_i\right] + \sum_{\mathcal{K}: i \notin \mathcal{K}} \frac{1-f}{1+f} p_i \log\left[\frac{1-f}{1+f} p_i\right] \right. \\ & \left. - \sum_{\mathcal{K}: i \in \mathcal{K}} \frac{1+f}{1-f} p_i \log[R] - \sum_{\mathcal{K}: i \notin \mathcal{K}} \frac{1-f}{1+f} p_i \log[R] \right\}; \end{aligned} \quad (20)$$

$$H_n(j) = - \sum_{K=0}^n \frac{(1-f)^K(1+f)^{n-K}}{2^n} p_j \sum_{\mathcal{K}} \log\left[\frac{p_j}{R}\right]. \quad (21)$$

The total posterior entropy $H_n = \sum_{i \text{ tested}} H_n(i) + \sum_{j \text{ not tested}} H_n(j)$. Analogous to Theorem 2a, we can show that this expression is minimal if $\sum_{i=1}^n p_i$ is maximal, or

equivalently, closest to $1/2$. The details are omitted here and can be obtained from the authors.

To see the second claim of the theorem, observe first that a larger f reduces the posterior entropy $H_n(i)$. In addition, a larger number of parallel tests decreases the posterior entropy convexely. We can show that $\frac{1}{2}[H_n(i) + H_{n+2}(i)] > H_{n+1}(i)$. Moreover, $[H_n(i) - H_{n+1}(i)]$ increases in f : a higher fidelity enhances the entropy reduction effect of a given number of tests. The proofs of these statements are messy and omitted here (they can be obtained from the authors).

We can write the optimal dynamic programming recursion, assuming the optimal policy of always testing the candidates with the largest probabilities, as: $V(H) = \min_n \{nc + c_\tau + V(H_n)\}$. As H_n decreases convexely in n , there is a unique minimum n^* . If we approximate H_n by a continuous function in n , the implicit function theorem implies

$$\frac{\partial n^*}{\partial f} = -\frac{\partial^2 H_n / (\partial f \partial n)}{\partial^2 H_n / \partial n^2} \geq 0.$$

n^* increases weakly in f because it is integer. This proves the second claim of the Theorem.

Proof of Theorem 3. We first calculate an upper limit on C_{mod} . As the M independent PEs can be tested in parallel, the costs of the tests simply add up. The time to test each PE is a random variable that can vary between 1 (first batch contains the solution) and $n(N)$ (last batch contains the solution). The expected time to test M PEs in parallel is the expectation of the maximum of these random variables. The expectation of the maximum of M independent uniformly distributed random variables is $\frac{M}{M+1}n$. From Corollary 1, the testing time distribution is skewed to the left: expected testing time is $(n(N) + 1)/3$. Thus, the expectation of the maximum is smaller than for a uniform distribution. The test costs simply add up for the M PEs. This gives the bound on the total cost in the middle column. The extreme cases for parallel and sequential testing (left and right columns) follow directly from Corollary 1.

For estimating C_{sequ} , assume first that the M PEs are tested sequentially, upstream before downstream. Then the total costs simply add up, both in time and in the number of tests, which gives the middle row of the Theorem. This is larger than C_{mod} for any n because $n/(M+1) < (n+1)/3$. It may be possible to reduce C_{sequ} by testing an upstream and a downstream PE in an overlapped manner. The best that can be achieved by overlapping is C_{mod} , provided that downstream picks the correct upstream alternative as the assumed solution and tests only its own alternatives compatible with this assumed upstream solution. In expectation, the overlapped cost is larger than C_{mod} because the assumed solution may be wrong, or downstream must test its own candidates in multiple versions corresponding to multiple upstream solutions. This proves the comparison statement in corollary 2.

Finally, we estimate C_{int} . In the integral case, the solution of one PE depends on the solutions of the others, and therefore, all combinations of alternatives must be tested. This is equivalent to one PE with N^M alternatives. This gives the third row of the Theorem. The conditions for the extreme cases (parallel or sequential testing) change because the number of alternatives is now different; a PE of N candidates may be tested sequentially, while it may be optimal to test partially in parallel in the PE of N^M candidates.

Inspection shows that for $2cN^M/c_\tau$ large, $C_{int} > C_{sequ}$. Numerical analyses show that $C_{int} > C_{sequ}$ for all possible parameter constellations as long as $3/8N \leq c/c_\tau$ holds (see Corollary 1). When delay costs are so high that this condition is not fulfilled, tests are performed in parallel (Corollary 1), and the total costs of testing multiple PEs become the same in both cases⁸. Again, C_{mod} is smallest, and $C_{int} > C_{sequ}$ iff $\frac{c}{c_\tau} > \frac{M-1}{N(N^{M-1}-M)}$. If c/c_τ is even smaller, it is optimal to test sequentially dependent PEs in parallel, incurring the extra cost of testing all combinations of alternatives in order to gain time. In this extreme case, $C_{int} = C_{sequ}$. This proves Theorem 3 and Corollary 2. $\square$

⁸Here, we assume that $3/8N^f > c/c_t$ also holds. If not, the integral design will not be tested fully in parallel, which makes the argument slightly more complicated (omitted here).

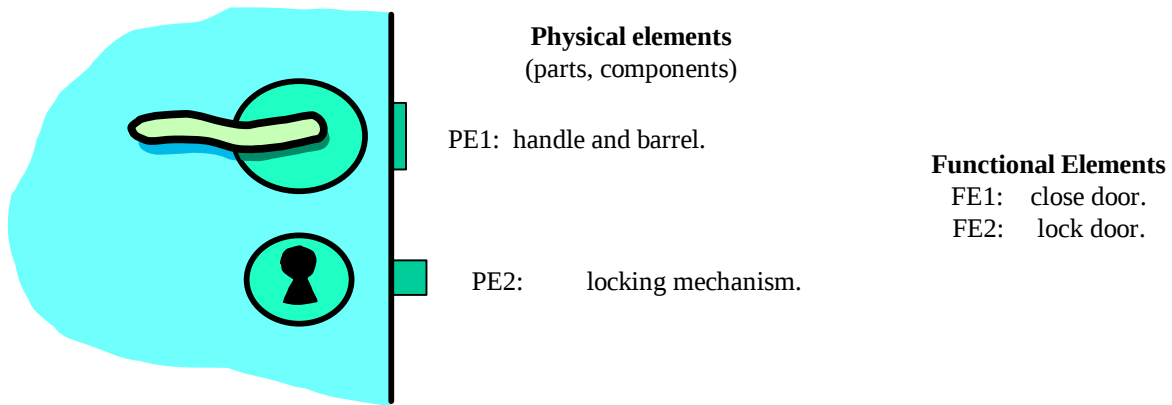


Figure 1: FEs and PEs in the Design of a Door

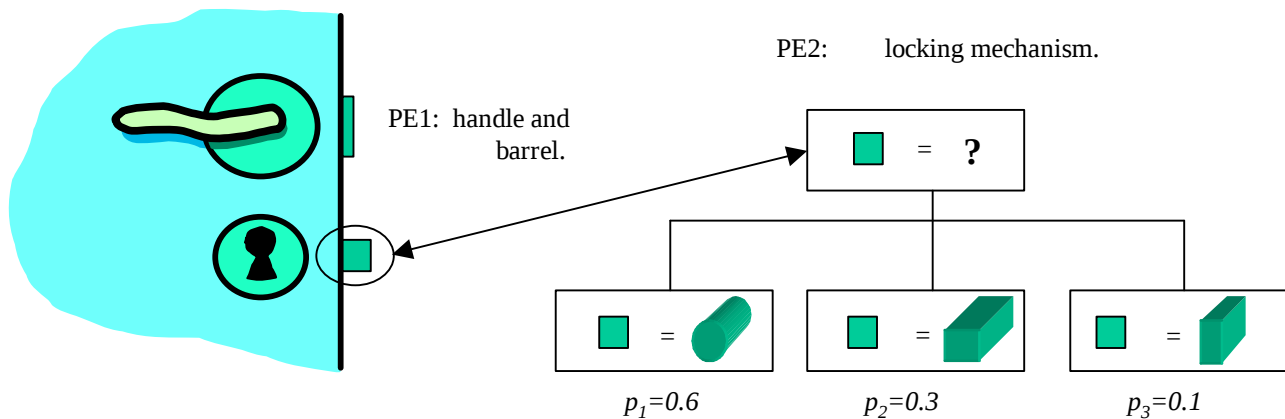


Figure 2: Solutions for the PE “locking mechanism” to fulfill the FE “lock the door”

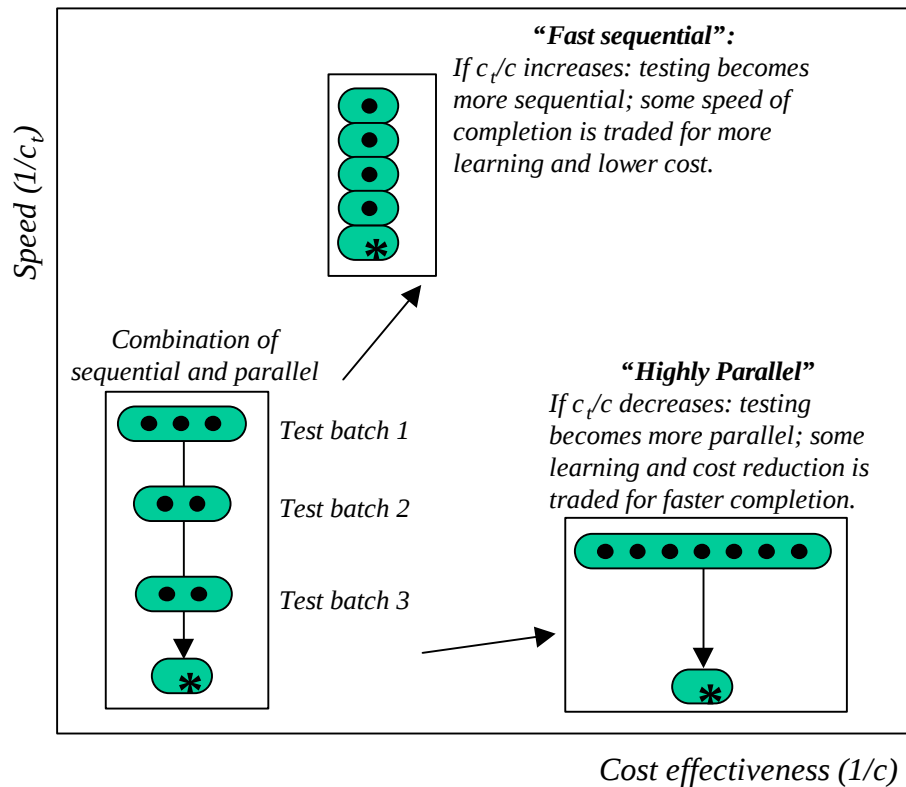


Figure 3: Impact of Test Speed and Cost on Testing Strategy

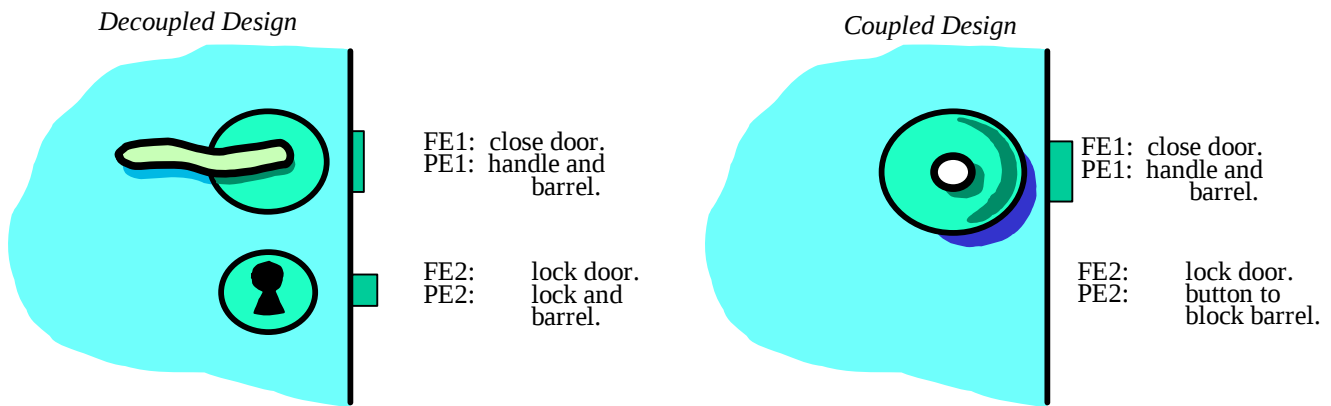


Figure 4: Functional Decoupling in the Design of a Door

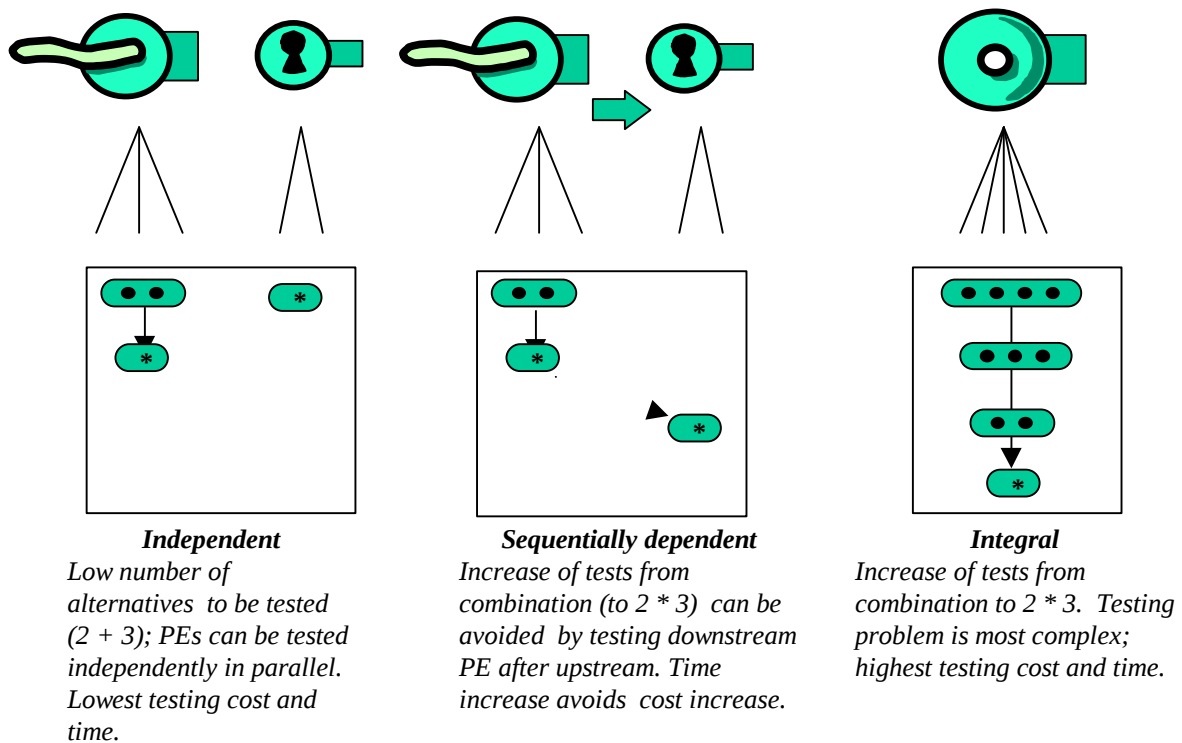


Figure 5: Impact of Architecture on Testing